

الجمهورية الجزائرية الديمقراطية
الشعبية

REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE

وزارة التعليم العالي والبحث
العلمي

Ministère de l'Enseignement Supérieur et de la Recherche Scientifique

جامعة سعيدة – د. الطاهر مولاي –

Université Saïda – Dr. Tahar Moulay –

Faculté des Mathématiques, Informatique et Télécommunications



MEMOIRE

Présenté pour l'obtention du **Diplôme de MASTER en Télécommunications**

Spécialité : Réseaux et Télécommunications

Par : M^{lle}. KAAROUR Kaouther

M^{lle}. KADA Rekia

Deep Learning pour le débruitage impulsif dans les systèmes OFDM

Soutenu, le 19 /06/ 2025, devant le jury composé de :

M. CHETIOUI Mohamed	MCA	Président
M. MOKADEM Djelloul	MCB	Rapporteur
M. GUENDOZ Mohamed	MCA	Examineur

2024 / 2025

Abstract

In recent years, Deep Learning (DL) has attracted increasing attention as a promising method for achieving noise cancellation. Unlike the traditional OFDM receiver, OFDM transmission is affected by non-Gaussian interference caused by impulsive phenomena.... A number of noise cancellation algorithms have been developed, based on the assumption that noise is Gaussian, such as additive white Gaussian noise (AWGN). It is within this framework that we propose a deep learning-based scheme to achieve impulsive noise cancellation to support the performance of traditional OFDM receivers. The application is specifically realized at the OFDM system level in a Rayleigh fading channel. Our simulation results clearly demonstrate the effectiveness of our solutions in cancelling impulsive noise and improving BER (bit error rate) performance, even in the presence of impulse noise interference.

Keywords: OFDM, Impulsive noise, Deep Learning (DL), BER.

Résumé

Ces dernières années, l'apprentissage profond (Deep Learning, DL) a attiré une attention de plus en plus importante en tant que méthode prometteuse pour réaliser l'annulation de bruit. À la différence du récepteur OFDM traditionnel, la transmission OFDM est affectée par des interférences non gaussiennes causées par des phénomènes impulsifs.. Ainsi, plusieurs algorithmes d'annulation de bruit ont été mis en place en reprenant l'hypothèse que le bruit est gaussien, à l'exemple du bruit blanc additif gaussien (AWGN). C'est dans ce cadre que nous proposons un schéma à base d'apprentissage profond pour réaliser l'annulation du bruit impulsif afin de soutenir les performances des récepteurs OFDM traditionnels. L'application est spécifiquement réalisée au niveau du système OFDM dans un canal d'évanouissement de Rayleigh. Les résultats de nos simulations montrent bien l'efficacité de nos solutions pour annuler le bruit impulsif et améliorer les performances BER (taux d'erreur binaire), même en cas d'interférence de bruit impulse.

Mots clés : OFDM, Bruit impulsif, Apprentissage en profondeur (DL), BER.

ملخص

في السنوات الأخيرة، اجتذب التعلم العميق (DL) اهتمامًا متزايدًا كطريقة واعدة لتحقيق إلغاء الضوضاء. على عكس مستقبل OFDM التقليدي، يتأثر إرسال OFDM بالتداخل غير الغاوسي الناتج عن الظواهر الاندفاعية. لذلك تم تطوير العديد من خوارزميات إلغاء الضوضاء بناءً على افتراض أن الضوضاء غاوسي، مثل الضوضاء البيضاء المضافة (AWGN) في هذا الإطار، نقترح مخططًا قائمًا على التعلم العميق لإجراء إلغاء الضوضاء الاندفاعية من أجل دعم أداء مستقبلات OFDM التقليدية. يتم تنفيذ التطبيق على وجه التحديد على مستوى نظام OFDM على مستوى نظام OFDM في قناة رايلي المتلاشية. تُظهر نتائج عمليات المحاكاة التي أجريناها بوضوح فعالية حلولنا في إلغاء الضوضاء الاندفاعية وتحسين أداء معدل خطأ البت (BER)، حتى في وجود تداخل الضوضاء الاندفاعية.

Dédicace

J'ai le plaisir de consacrer ce petit travail :

Au meilleur papa (Abdelkrim) :

Tu m'as toujours montré la bonne voie. Merci pour tes valeurs nobles, ton éducation et ton soutien constant.

A ma très chère maman (Karima) :

Tu incarnes la bonté absolue, la tendresse infinie et le dévouement sans faille. Ton soutien et tes prières ont toujours été mon refuge. Que, par moi, tu ressenties cette fierté tant méritée. Qu'Allah, le Tout-Puissant, veille sur toi, te comble de santé et de bonheur.

A mon Cher frère (Farouk) que Dieu le protège, une longue vie Inchallah et que je t'aime bien.

J'adresse aussi mes dédicaces à mes amies
(Daisy, Rania, Rouka, Rajouaa, Asmaa).

KAOUTHER

Dédicace

.

A l'homme, mon précieux offre du Dieu, qui doit ma vie, ma réussite et tout mon respect: mon cher père MOKHTAR

A la femme qui a souffert sans me laisser souffrir, qui n'a jamais dit non âmes exigences et qui n'a épargné aucun effort pour me rendre heureuse: mon adorable mère NORA

À mon frère YAHIA ma sœur IMANE

À toutes ma famille de près ou de loin

À ma binôme KAOUTHER tu es bien plus coéquipière merci d'avoir été là à chaque étape de ce travail

REKIA

Remerciements

Nous remercions d'abord tout d'abord ALLAH le tout puissant de nous avoir
donné la

santé, la patience, la puissance et la volonté pour accomplir ce mémoire.

Nous exprimons, notre profonde gratitude et tout notre amour à
nos parents, nos sœurs et frères, qui ont su nous faire confiance et
nous soutenir en toutes circonstances,

Nous tenons en particulier à remercier notre promoteur,

Dr Mokadem Djelloul

pour avoir accepté la charge d'être

Rapporteur de ce mémoire, nous le remercions pour sa disponibilité, ses
pertinents conseils sa patience et pour les efforts qu'il a consenti durant la
réalisation de ce mémoire. Qu'il trouve ici l'expression de notre
reconnaissance

et de notre respect.

Nos vifs remerciements doivent également accompagner membres du jury
pour honorer notre soutenance de Master Merci pour votre présence.

Ainsi qu'à tous nos proches amis qui nous ont toujours soutenus et
encouragés même dans les périodes les plus difficiles.

« Merci »

KAOUTEHER & REKIA

Sommaire ou tables de matières

Les résumés	Erreur ! Signet non défini.
Dédicace	3
Remercîments.....	5
Sommaire ou tables de matières.....	6
Liste des tableaux	9
Listes des figures.....	10
Acronymes	11
Introduction générale.....	2
CHAPITRE I. L'évolution des réseaux mobiles de la 1G vers la 5G	5
I.1 Introduction	5
I.2 Evolution des systèmes radio cellulaires :	5
I 2.1 Les réseaux mobiles de première génération (1G) :	5
I 2.2 Les réseaux mobiles de deuxième génération (2G) :	6
I.1.2.1 Le réseau GSM :	6
I.1.2.2 Architecture GSM	7
I 2.3 Les réseaux mobiles de troisième génération (3G)	8
I.1.3.1 UMTS	8
I.1.3.2 Description générale de l'architecture d'un réseau UMTS :.....	8
I.1.3.3 Le HSPA :	9
I 2.4 Les réseaux mobiles quatrièmes générations (4 G) :	10
I.1.4.1 De l'UMTS au LTE :	10
I.1.4.2 Architecture du réseau LTE :	11
I.1.4.3 Le réseau LTE-A :	12
I 2.5 Les réseaux mobiles de cinquième génération (5G) :	12
I.1.5.1 Réseau cœur 5G (5GC) :	14
I.3 Comparaison entre les 5 générations de communication :	15
I.4 Conclusion :	17
CHAPITRE II. OFDM	19
II.1 Introduction :	19
II.2 Principe de fonctionnement :	19
II.3 Chaîne de transmission :	20
II 3.1 Emetteur :	20
II 3.2 Canal de transmission :	22
II.1.2.1 Canal de propagation par trajets multiples :	24

II.1.2.2 Canal à Bruit blanc additif gaussien (AWGN).....	25
II.1.2.3 Canal de propagation à bruit impulsif	26
II 3.3 Récepteur :	30
II 3.4 Conclusion :	31

CHAPITRE III. DEEP LEARNING POUR LA REDUCTION DU BRUIT DANS UNE

CHAINE OFDM	33
III.1 Introduction :	33
III.2 Les réseaux neuronaux :	33
III 2.1 Présentation sommaire du neurone biologique :.....	33
III 2.2 Les réseaux neuronaux artificiels :	34
III 2.3 Modèle mathématique d'un neurone artificiel :	36
III.3 Apprentissage des réseaux neuronaux artificiels :	36
III 3.1 Vocabulaire des Réseaux Neuronaux :	37
III.4 -Les réseaux CNN (Convolutional Neural Networks)	37
III.5 -Architecture d'un Convolutional Neural Network-CNN	37
III.6 L'auto encodeur	38
III 6.1 L'architecture d'un auto-encodeur :	39
III 6.2 Paramètres d'entraînement pour un auto-encodeur :	40
III 6.3 Principe de fonctionnement d'un auto-encodeur :	40
III.7 DENOISING AUTO-ENCODER (DAE) DANS UNE CHAINE OFDM	42
III 7.1 . Génération de signaux bruités – Partie émission :	42
III 7.2 Modélisation du bruit.....	42

CHAPITRE IV. Résultats de simulation44

IV.1 Introduction.....	44
IV.2 Présentation de Google colab	44
IV.3 Pourquoi choisir Google Colab ?	44
IV.4 Simulation et Résultats.....	45
IV 4.1 Simulation de la chaine OFDM	45
IV.5 PARTIE 01 : simulation avec bruit Gaussien	46
IV 5.1 Analyse de la Densité Spectrale de Puissance (DSP) et du Taux d'Erreur Binaire (BER) pour les Modulations BPSK, QPSK et 16-QAM en présence de Bruit Gaussien	46
IV 5.2 Conclusion	50
IV.6 Partie 02 : simulation avec bruit impulsif.....	50
IV 6.1 . COMMENT IMPLEMENTER UN AUTO ENCODEUR POUR LE DEBRUITAGE D'UN SIGNAL (DENOISING AUTO-ENCODEUR DAE)	50
IV 6.2 Points importants à considérer pour votre application spécifique (par exemple, le débruitage de signaux OFDM) :	51

IV 6.3 Implémentation de l'auto-encodeur dans la chaîne OFDM.....	52
IV 6.4 . L'Émetteur (Bloc Bleu) : Préparation du Signal pour la Transmission	55
IV 6.5 Le Canal (Bloc Rouge) : Introduction des Perturbations	56
IV 6.6 Le Récepteur (Bloc Vert) : Récupération du Signal Original	56
IV 6.7 CONCLUSION.....	Erreur ! Signet non défini.
IV.7 Entraînement de l'Auto-Encodeur :.....	57
IV.8 Utilisation pour la Transmission :	57
IV.9 Conclusion	58
Conclusion générale	59
References bibliographiques	62

Liste des tableaux

Tableau 1 : Les différents standards du réseau [1].....	6
Tableau I-2 : Comparaison entre les 5 générations.....	17
Tableau II-1 : les équations des cas particuliers de α - stable	30
Tableau III-1 : Analogie entre neurones biologiques et artificiels.....	35
Tableau IV-1 Paramètres de simulation	45

Listes des figures

Figure I-1 : Évolution de communication mobile 1G à 5G	5
Figure I-2: Architecture GSM [4].....	7
Figure I-3 : Architecture du réseau UMTS	9
Figure I-4 : Architecture du réseau LTE [9]	11
Figure I-5 : Architecture de la cinquième génération (5G).....	13
Figure I-6 : Architecture du NG-RAN.....	14
Figure I-7 : Architecture du 5G Cœur [6].	14
Figure II-1 : Représentations fréquentielle et temporelle du signal OFDM	19
Figure II-2 : Diagramme en bloc de la chaine de transmission OFDM.....	20
Figure II-3 représentation temporelle et vectorielle [21].....	21
Figure II-4 : propagation par trajets multiples.....	23
Figure II-5 : Bruit pour différentes valeurs de A et $\Gamma=0.001$	27
Figure II-6 : (a) deux modèles de bruit de classe d'état et (b) modèle de bruit Bernoulli-gaussien ..	29
Figure III-1 : Diagramme de Venn	33
Figure III-2 : Schéma d'un neurone biologique [38]	34
Figure III-3 : Neurone biologique	34
Figure III-4 : l'architecture d'un Réseau de neurones	35
Figure III-5 Modèle d'un neurone artificiel (Marc Parizeau., 2004)	36
Figure III-6 : l'architecture d'un CNN	38
Figure IV-1 : the collaboriez	44
Figure IV-2 Schéma bloc de la chaine de transmission.....	45
Figure IV-3 : DSP de BPSK avec bruit Gaussien	47
Figure IV-4 : DSP de QPSK avec bruit Gaussien	47
Figure IV-5 : DSP de 16-QAM avec bruit Gaussien	48
Figure IV-6 : la comparaison des densités spectrales de puissance avec bruit Gaussien.....	48
Figure IV-7 : la courbe de BER vs SNR avec bruit gaussien.	49
Figure IV-8 : la comparaison BER simulé s théorique avec bruit Gaussien	50
Figure IV-9 : Débruitage Auto-encodeur en action	51
Figure IV-10 : Schéma fonctionnel.....	53
Figure IV-11 : Schéma fonctionnel d'un système de communication OFDM basé sur un Auto-	55

Acronymes

1G	Première Génération.	D	
2G	Deuxième Génération.	DAE	Denoising Autoencoder.
3G	Troisième Génération	DCS	Digital Cellular System
4G	Quatrième Génération.	E	
5G	Cinquième Génération	EDGE	Enhanced Data Rates for GSM Evolution.
3GPP	3rd Generation Partnership Project	EIR	Equipment Identity Registration
A		EPC	Evolved PacketCore Network.
AE	Auto Encodeur	eNB	Evolved NodeB.
AMPS	Advanced Mobile Phone System.	F	
AuC	Authentication Center	FDMA	Frequency Division Multiple Access.
AMF	Access Mobile Function.	FTT	Fourier Transform.
AF	Application Function.	G	
AWGN	Additif White Gaussian Noise.	GSM	Global system for mobile communication.
B		gNB	Next Generation NodeB.
BER	BER : Bit Error Rate.	gNB-CU	gNB-Central Unit.
BPSK	Binary Phase Shift Keying.	gNB-CU-CP	gNB Central Unit - Control Plane.
BS	Base Station.	gNB-CU-UP	gNodeB Central Unit - User Plane.
BSS	Base Station Sub system.	GPRS	General Packet Radio Service.
BSC	Base Station Controller.	H	
BTS	Base Transceiver Station.	HLR	Home Location Register.
C		HSS	Home Subscriber Server.
CDMA.	Code Division Multiple Access.	HTTP	Hypertext Transfer Protocol.
CP.	commutation paquets.		
CN	Core Network.		
CNN.	Convolutional Neural Network.		
CUPS	Control and User Plane Separation.		

I

IFF	Inverse Fast Fourier
IM	International Mobile

L

LO	Line-Of-Sight
LTE	Long Term Evolution.
LTE	LTE Advanced.

M

M	Multimedia Messaging
MS	Mobile Station International

N

NL	None Line-Of-Sight.
N	Nordic Mobile Téléphone.
NG	Next Generation Radio

NF	Network Function.
-----------	-------------------

O

OF	Orthogonal Frequency
-----------	----------------------

P

PD	Personal Digital
PD	Probability Density Function.
PG	Packet Data Network
PC	Policy and Charging Rules
PD	Packet Data Network.

Q

QP	Quadrature Phase Shift
QA	Quadrature Amplitude

R

Re	Rectified Linear Unit.
-----------	------------------------

RN	. Recurrent Neural Network.
-----------	-----------------------------

S

	System Architecture
SG	Serving Gateway.
SM	Short Message Service
SN	Signal to Noise Ration.

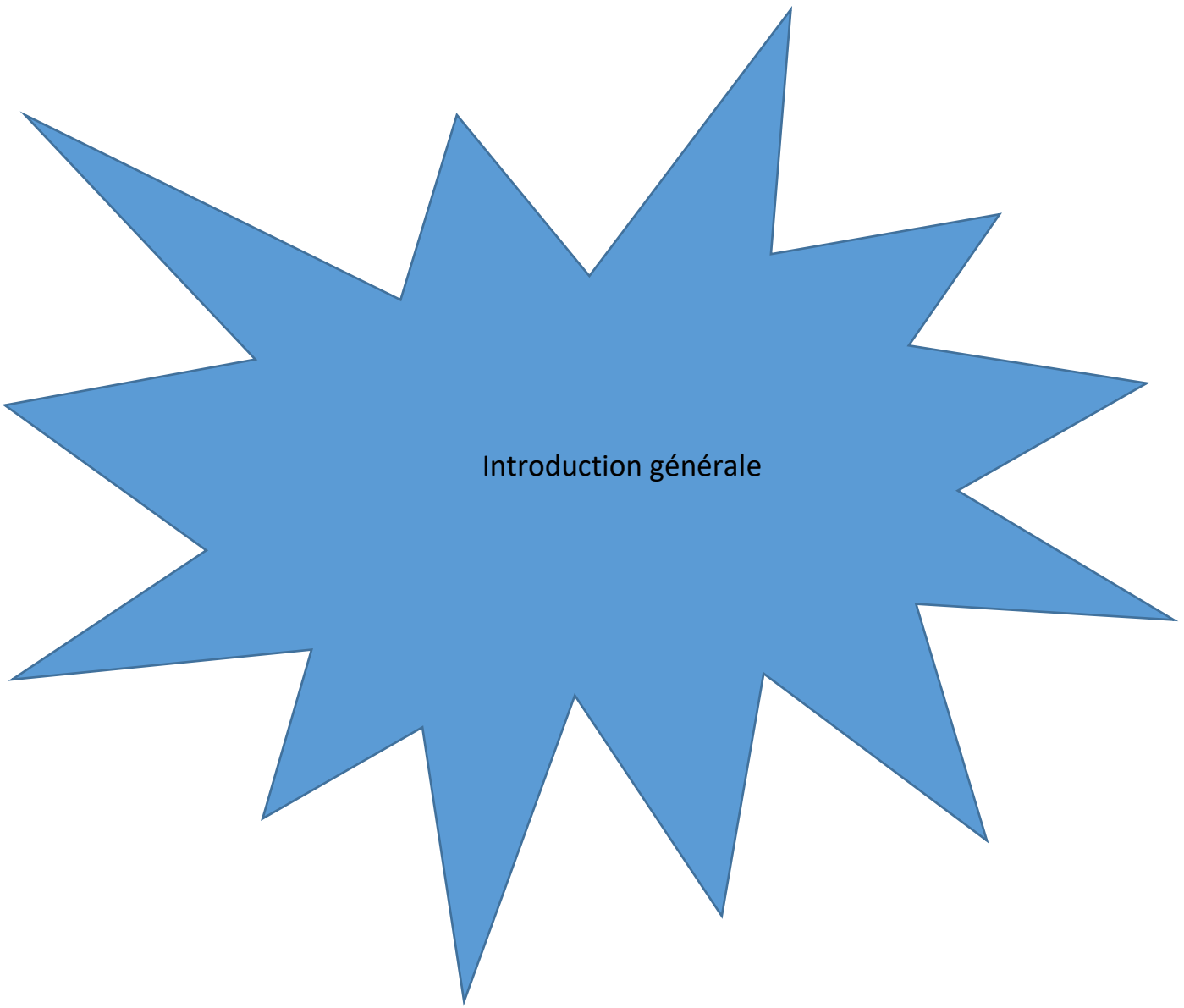
T

TE	Taux Erreur Binaire.
TA	Total Access
TD	Time Division Multiple
TM	Temporary Mobile

U

UE	User Equipment.
U	Universal Mobile

V



Introduction générale

Introduction générale

Au fil du temps, les systèmes de communication sans fil ont connu une évolution remarquable grâce aux avancées technologiques dans des disciplines variées comme le traitement du signal, la micro et nanoélectronique, le développement logiciel et les mathématiques. Ces progrès ont conduit à l'émergence successive des différentes générations de téléphonie mobile. Depuis les années 1980, où le premier système de radio mobile analogique (1G) a vu le jour, jusqu'à l'introduction de la deuxième génération (2G) qui marqua la transition vers le numérique, chaque étape a révolutionné notre manière de communiquer. La troisième génération (3G) a ensuite ouvert la voie aux communications à haut débit, avant que la quatrième génération (4G), désormais couramment utilisée, ne nous offre des connexions rapides et stables. Enfin, la cinquième génération (5G), en cours de déploiement, promet d'amplifier encore davantage nos capacités de communication.

Pour répondre aux besoins croissants des utilisateurs en termes de débit et de mobilité, les modulations mono-porteuses traditionnelles se révèlent insuffisantes, notamment face aux défis posés par la sélectivité fréquentielle des canaux et les interférences dues aux trajets multiples. En conséquence, les réseaux sans fil ont opté pour la technique OFDM (Orthogonal Frequency Division Multiplexing), une méthode de modulation multi-porteuse qui a grandement renforcé leur capacité de transmission.

Les canaux de transmission sans fil ont évolué en termes de caractéristiques et de nature de bruit. Au début, les modèles de bruits utilisés étaient simples et se limitaient au bruit blanc additif gaussien (AWGN), qui est un bruit aléatoire et stationnaire avec une densité spectrale de puissance constante dans toutes les bandes de fréquences. Cependant, avec le développement des technologies de communication sans fil, les canaux de transmission sont devenus plus complexes et ont commencé à présenter des caractéristiques de bruit impulsif, qui est un bruit non stationnaire et non gaussien avec une densité spectrale de puissance variable dans différentes bandes de fréquences.

Ces dernières années, l'apprentissage profond est devenu une méthode populaire pour résoudre les problèmes de traitement du signal, y compris l'annulation du bruit. Cette méthode utilise des réseaux de neurones artificiels pour apprendre à annuler le bruit de manière automatique. Et c'est dans ce cadre s'articule notre projet de fin d'études dans lequel nous allons utiliser les techniques de Deep learning, à savoir les réseaux de neurones CNN (Convolutional Neural Network) pour supprimer le bruit impulsif des signaux OFDM. Nous avons organisé notre mémoire de fin d'études en quatre chapitres.

Dans le premier chapitre, nous présentons une vue d'ensemble des différentes générations de téléphonie mobile, de la 1G à la 5G, en expliquant les normes de téléphonie mobile et les technologies utilisées dans chaque génération.

Dans le deuxième chapitre, nous décrirons en détail la technologie OFDM (Orthogonal Frequency Division Multiplexing), son principe de fonctionnement, et les différents blocs constituant la chaîne de transmission. En plus nous allons aborder les différents types de canaux de communication. Nous commençons par une explication détaillée du canal à bruit blanc additif gaussien (AWGN), suivi du canal de

propagation à trajets multiples avec les distributions Rayleigh, et enfin en nous allons simuler un bruit impulsif dans un canal de propagation. Dans le troisième chapitre, nous allons explorer les concepts clés de l'apprentissage profond, en commençant par une présentation sommaire de réseaux neuronaux artificiels. Enfin, nous allons aborder un concept important de l'apprentissage profond, l'auto-encodeur, qui sera utilisé par la suite dans l'élimination de bruit impulsif.

Et enfin notre mémoire sera clôturé par une conclusion générale dans laquelle, nous allons présenter de façon sommaire les résultats obtenus ainsi que les défis que nous allons rencontrer.

Chapitre I : L'évolution des
réseaux mobiles de la 1G vers la 5G

I.1 Introduction

Dans ce chapitre, nous expliquerons l'évolution des technologies mobiles au fil des années. Il existe déjà des réseaux mobiles de 1G, 2G, 3G, 4G et maintenant 5G. L'ensemble de ce processus d'évolution a pris environ 40 ans. En raison de l'énorme besoin de connexions à travers le monde, les normes de communication mobile ont considérablement progressé dans leurs performances et de leur sécurité pour prendre en charge d'avantage d'utilisateurs [1].

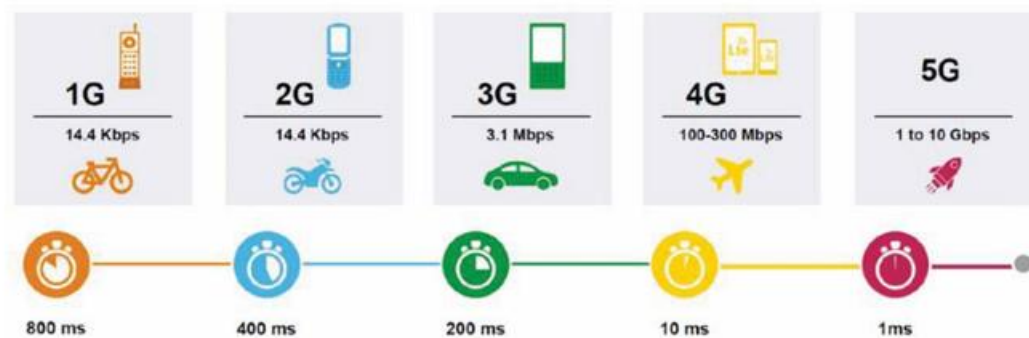


Figure I-1 : Évolution de communication mobile 1G à 5G

I.2 Evolution des systèmes radio cellulaires :

Les réseaux mobiles de première génération (1G) :

Il y a de cela un certain temps que le web explore de nouvelles dimensions, et rend l'univers technique incomparablement diversifié. Faisons un bond dans le passé à 1G, le point initial. 1G est la première génération de communications mobiles, apparaissant dans les années 1980. Les systèmes cellulaires utilisés étaient des systèmes à commutation de circuit et analogiques conçus pour offrir des services de téléphonie, où chaque appel nécessitait un canal de fréquence étroite [2]. Il y a eu plusieurs systèmes de téléphonie mobile qui ont été mis en œuvre dans certaines parties du monde comme le NMT, AMPS et TACS.

Standard	utilisation	Technologie
Système avancé de téléphonie mobile(AMPS)	Amérique du Nord, 1980	Cellulaire, analogique,FDMA
Système de communication à accès total (TACS)	UK, 1983	Cellulaire, analogique,FDMA
Téléphone mobile nordique(NMT)	Suède, Russie, 1 octobre 1981	Cellulaire, analogique,FDMA
Le réseau radiotéléphonique C	Allemagne, 1985	Cellulaire, analogique, FDMA, Itinérance.

Tableau 1 : Les différents standards du réseau [1]

Les réseaux mobiles de deuxième génération (2G) :

(2G) ou les réseaux mobiles de deuxième génération ont été développés autour des années 1990, marquant le passage de l'analogique au numérique. Le développement des semi-conducteurs a permis d'intégrer plus d'électronique à moindre coût dans des formats plus petits. Cela signifie que les téléphones mobiles 2G sont devenus abordables pour le grand public. En plus de la téléphonie vocale, cette nouvelle génération offrait d'autres services tels que le SMS, le MMS, le télécopie, et l'accès à des réseaux numériques comme l'internet et l'ISDN (réseau numérique à services intégrés). Ces innovations ont élargi les usages de la téléphonie mobile au-delà des capacités d'appels vocaux.

Il existe principalement trois normes dans le système de communication mobile ; ce sont le GSM (Système mondial de communication mobile), le 2.5G GPRS (Service général de paquet radio) et l'EDGE ou 2.75G (Taux de données amélioré pour l'évolution du GSM).

I.1.1.1 Le réseau GSM :

Le GSM, ou Système Mondial de Communication Mobile, a été lancé en 1990 par la CEPT. C'est la première norme numérique pour téléphones portables. Il a remplacé les anciens réseaux analogiques. Il est vite devenu la technologie principale. Il offrait une meilleure qualité sonore, les SMS et une couverture partout dans le monde.

C'est une technologie de communication très répandue. Elle utilise des canaux radio numériques pour envoyer la voix et les données des téléphones portables. Selon le système, différentes fréquences sont utilisées. Par exemple, les fréquences 450 Mhz pour les SMS, 900 Mhz pour le GSM et 1800 Mhz pour le DCS 1800. DCS signifie Système Cellulaire Numérique.

En Europe, on peut trouver ces deux types de réseaux [3]: Chapitre I L'évolution des réseaux mobiles de la 1G vers la 5G 7 → GSM 900 : La bande 890-915MHz pour l'Up Link et la bande 935-960 MHz pour le Down Link. → Le GSM 1800 : également appelé DCS 1800, la bande Up Link 1710MHz-1785MHz et La bande Down Link 1805MHz-1880MHz. Les dispositifs qui opèrent à la fois sur les fréquences 900 et 1800 sont connus sous le nom appelés GSM dual band ou simplement dual band.

I.1.1.2 Architecture GSM

Pour offrir un service accessible, les opérateurs de réseaux de téléphonie mobile installent un certain nombre de stations de base (appelées BS : stations de base) dans la zone qu'ils veulent couvrir. Le but est que votre téléphone reste toujours à quelques kilomètres d'une de ces stations. La zone où un téléphone peut se connecter à une station de base s'appelle une cellule. C'est pourquoi l'espace de couverture est divisé en cellules voisines. On appelle une station mobile tout terminal capable de communiquer sur un réseau.

Un téléphone mobile se compose essentiellement d'un émetteur-récepteur et d'une logique de commande [3].

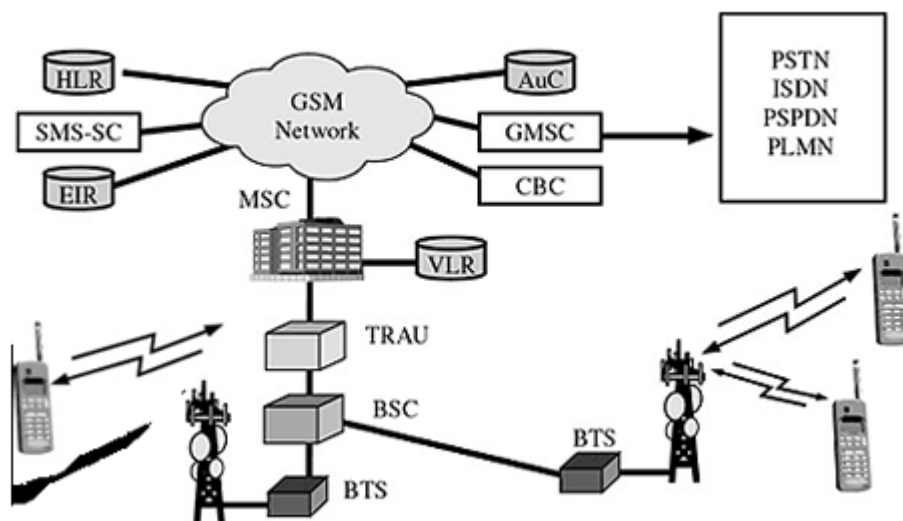


Figure I-2: Architecture GSM [4]

Le sous-système BSS (Base Station Sub system) : Le BSS (Base Station Subsystem) gère la qualité des liaisons radio dans un réseau GSM, comprenant le BTS (station de base) et le BSC (contrôleur de station de base).

- ❖ BTS (Base Transceiver Station) : Gère la communication radio avec les téléphones mobiles dans une cellule donnée, en assurant la transmission et la réception des signaux.
- ❖ BSC (Base Station Controller) : Supervise plusieurs BTS, gère les ressources radio, le handover entre cellules et l'acheminement des appels vers le réseau central.

Le sous-système NSS (Network Subsystem) : Le NSS (Network Subsystem) gère l'ensemble des fonctions de commutation et de gestion des abonnés dans un réseau GSM. Il comprend plusieurs composants clés :

- ❖ MSC (Mobile Switching Center) : Assure la commutation des appels et la gestion des communications avec d'autres réseaux mobiles et terrestres. Il inclut le sous-système de commutation et de contrôle.

- ❖ HLR (Home Location Register) : Base de données principale stockant les informations des abonnés, comme le MSISDN (numéro de téléphone) et le type de service.
- ❖ VLR (Visitor Location Register) : Base de données temporaire qui stocke les informations des abonnés se trouvant dans une zone différente de leur zone d'enregistrement (une sorte de copie locale du HLR pour un usage temporaire).
- ❖ EIR (Equipment Identity Register) : Contient la base de données des IMEI (numéros uniques attribués à chaque appareil mobile) pour identifier et suivre les équipements.
- ❖ AuC (Authentication Center) : Gère l'authentification des utilisateurs pour l'accès au réseau mobile, le cryptage des communications sans fil et l'attribution d'identités temporaires TMSI (Temporary Mobile Subscriber Identity)

Les réseaux mobiles de troisième génération (3G)

La troisième génération 3G de réseaux cellulaires fait référence à la technologie utilisée dans les télécommunications mobiles pour permettre le transfert de données à haut débit, les services multimédias sophistiqués comme le streaming vidéo, les jeux vidéo, la navigation GPS... [5]. Au début des années 2000, cette technologie a été introduite pour la première fois, remplaçant les anciens systèmes de seconde génération. Elle permettait d'accéder à Internet plus rapidement et de supporter un grand nombre d'utilisateurs en même temps, grâce à des technologies avancées comme le CDMA (Code Division Multiple Access).

I.1.1.3 UMTS

L'UMTS (Système Universel de Télécommunications Mobiles) est une technologie de communication mobile qui fait partie de la famille des réseaux 3G, conçue dans le cadre des IMT-2000 et basée sur la norme mondiale GSM. Même si l'UMTS offre plus de fonctionnalités, il n'a pas été créé pour remplacer le GSM, mais plutôt pour venir le compléter. Cela nous permet d'utiliser les deux technologies sur le même téléphone. L'UMTS permet un transfert de données plus rapide, offre une meilleure couverture cellulaire, dispose d'une large bande passante et utilise efficacement le spectre radio.

Bien que l'UMTS utilise la technologie d'accès multiple par répartition en code (CDMA), il a une bande passante plus large que les autres systèmes CDMA par exemple CDMA2000 d'où l'appellation Widebande CDMA [6].

I.1.1.4 Description générale de l'architecture d'un réseau UMTS :

L'architecture d'un réseau UMTS est constituée de plusieurs entités physiques regroupées en domaines selon leur rôle au sein du réseau. L'architecture d'un réseau UMTS est composée de trois domaines (Figure I.3). Les domaines UTRAN et CN font partie du domaine de l'infrastructure du réseau [5].

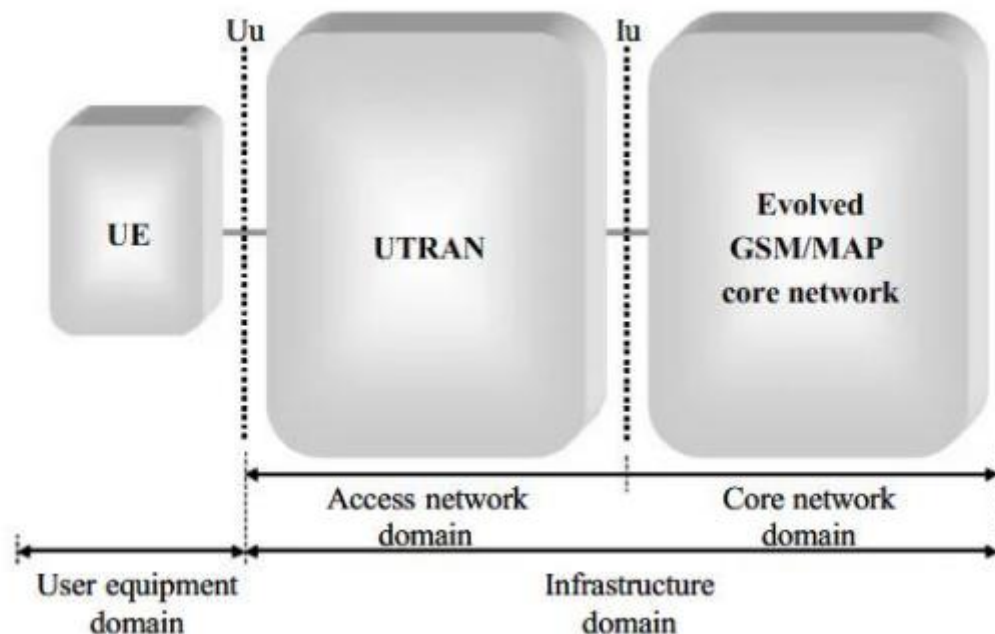


Figure I-3 : Architecture du réseau UMTS

- **UE (User Equipment)** : C'est le terminal mobile dans un réseau UMTS (similaire à une station mobile dans le réseau GSM/GPRS). L'UE permet à un abonné d'accéder aux services UMTS via l'interface radio Uu.
- **UTRAN (Universal Terrestrial Radio Access Network)** : C'est l'infrastructure qui permet à l'UE d'accéder au réseau central CN (Core Network). L'UTRAN gère les ressources radio, contrôle l'admission des connexions, et établit les bearers (canaux de communication) pour que l'UE puisse échanger des données avec le CN.
- **CN (Core Network)** : C'est la partie du réseau UMTS chargée de gérer les services de télécommunication pour chaque abonné. Cela inclut la gestion des appels (voix et données), l'authentification de l'UE, l'interconnexion avec les réseaux externes (mobiles et fixes), et la facturation des utilisateurs. Le CN gère l'établissement, la terminaison et la modification des appels, en utilisant des mécanismes de commutation par circuits ou par paquets.

I.1.1.5 Le HSPA :

Le High-Speed Packet Access (HSPA), aussi appelé 3G+, est une technologie conçue pour rendre la connexion internet sur les réseaux UMTS plus rapide et plus fiable. Elle se compose principalement de deux parties : le HSDPA (High-Speed Downlink Packet Access), qui accélère la vitesse de téléchargement des données, et le HSUPA (High-Speed Uplink Packet Access), qui rend le téléversement plus rapide. Cependant, même avec l'arrivée du HSPA, l'évolution d'UMTS ne s'arrêtait pas là. Le HSPA+ (ou Evolved HSPA) a permis d'aller encore plus loin en proposant des débits plus élevés et une meilleure gestion des ressources du réseau, renforçant ainsi les performances des réseaux 3G.

Les réseaux mobiles quatrièmes générations (4 G) :

La 4G, ou réseau de quatrième génération, offre des débits de données pouvant atteindre 20 Mbps tout en assurant une bonne qualité de service. L'un de ses grands avantages est sa capacité à gérer beaucoup de trafic, ce qui la rend parfaite pour les zones très peuplées. Elle optimise aussi l'utilisation du spectre en donnant la priorité au trafic selon le type d'application, ce qui lui permet de s'adapter rapidement aux différentes demandes. La 4G a été conçue pour offrir des services mobiles à large bande, beaucoup plus rapides que la 3G, et facilite la diffusion de contenus multimédias de haute qualité. Cela permet de supporter une variété d'applications, comme le streaming vidéo, les jeux en ligne ou la visioconférence. En plus, la 4G offre une plus grande largeur de bande aux appareils qui se déplacent rapidement, comme ceux dans les véhicules en mouvement, même dans des zones où le réseau est faible.

I.1.1.6 De l'UMTS au LTE :

En 2004, le 3GPP a commencé à étudier comment faire évoluer l'UMTS sur le long terme. Leur but était de garantir que les systèmes de communication mobile 3G restent compétitifs pendant plus de dix ans, en offrant des débits de données plus rapides et des temps de réponse plus faibles pour répondre aux besoins futurs des utilisateurs. Dans cette nouvelle architecture, le cœur de réseau amélioré basé sur le principe de paquets (EPC) remplace directement la partie à commutation de circuits de l'UMTS et du GSM. Il n'y a pas d'équivalent pour la commutation de circuits, ce qui permet au LTE d'être parfaitement adapté à la transmission de données, mais cela signifie aussi que les appels vocaux doivent être traités par d'autres techniques. Le réseau radio terrestre amélioré (E-UTRAN) gère les communications radio entre le réseau et le mobile, ce qui remplace directement l'UTRAN. Le mobile, souvent appelé l'équipement de l'utilisateur, fonctionne de manière très différente de ce qu'il faisait avant. Cette nouvelle architecture a été conçue dans le cadre de deux grands projets du 3GPP : l'évolution de l'architecture du système (SAE), qui couvre le réseau central, et l'évolution à long terme (LTE), qui concerne le réseau radio d'accès, l'interface aérienne et la téléphonie mobile.

L'ensemble du système est officiellement connu sous le nom de système de paquets évolué (EPS), alors que le sigle LTE fait uniquement référence à l'évolution de l'interface radio [8].

I.1.1.7 Architecture du réseau LTE :

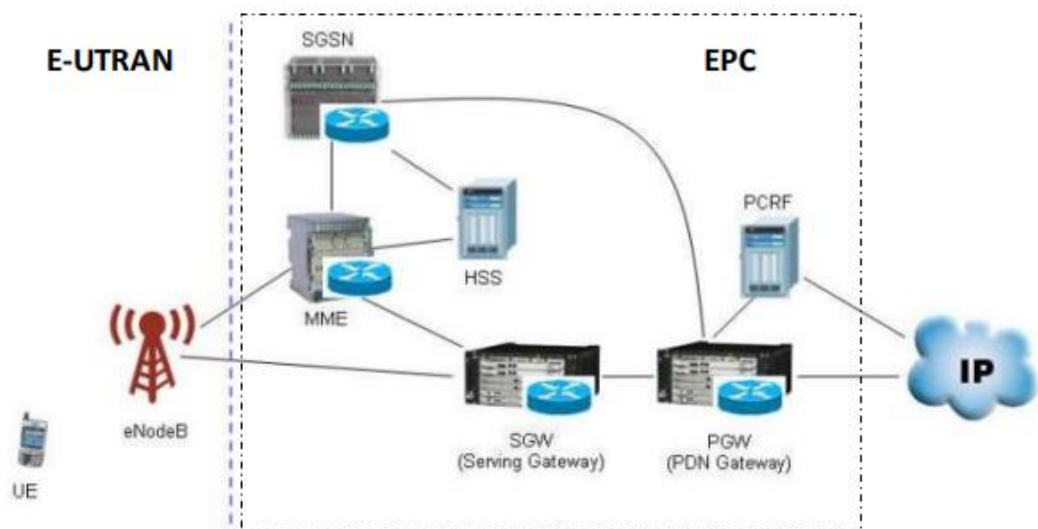


Figure I-4 : Architecture du réseau LTE [9]

▪ Le réseau d'accès E-UTRAN :

L'E-UTRAN, ou Réseau d'Accès Radio Terrestre Universel Évolué, c'est la partie radio du réseau LTE. En gros, c'est ce qui relie ton téléphone ou autre appareil mobile au réseau, en passant par des stations appelées Node B. Elle gère aussi la façon dont les ressources radio sont utilisées et communique avec le cœur du réseau pour que tout fonctionne bien.

Le Node B (evolved Node B) est la station de base sans fil utilisée dans les réseaux mobiles LTE. Elle joue un rôle clé en assurant la connexion des appareils mobiles au réseau LTE. Le Node B est responsable de plusieurs fonctions cruciales :

- Gestion de la connexion entre les UE et le réseau.
- Compression et décompression des en-têtes de données pour optimiser l'utilisation du réseau.
- Contrôle d'accès aux ressources radio.
- Ordonnancement des ressources pour assurer une transmission de données fluide et efficace.
- Sécurité des communications via des mécanismes de cryptage et d'authentification.
- Routage des données vers le réseau central pour le traitement et la transmission finale.

▪ Evolved Packet Core :

Le cœur d'un réseau LTE est représenté par l'EPC (Evolved Packet Core). L'EPC est composé de plusieurs nœuds, dont les principaux sont MME, SGW, PGW et HSS. Ces nœuds offrent de diverses fonctionnalités telles que la gestion de la mobilité, l'authentification, la gestion des sessions, l'établissement de support de communication et l'application de différentes qualités de service. L'EPC a une architecture IP plate qui permet au réseau de gérer une grande quantité de trafic de données de manière efficace et rentable [10]

- L'entité MME signifie "Mobility Management Entity". La MME est responsable de l'authentification de l'UE. En outre, elle suit la localisation de l'UE. Il sélectionne également le SGW et le PGW appropriés qui doivent desservir cet UE.
- L'entité SGW signifie "Serving Gateway" (passerelle de desserte). Afin d'éliminer tout effet sur les données de l'utilisateur lorsque l'UE se déplace entre différents eNodeB, le SGW fonctionne comme un point d'ancrage pour les données de l'utilisateur, lorsque l'UE se déplace entre différents eNodeB. En outre, le SGW transmet les données de l'utilisateur entre l'eNB et le PGW.
- L'entité PGW désigne la "Packet Data Network Gateway". Le PGW est le nœud qui relie le réseau LTE au PDN.
- L'entité HSS désigne le "Home Subscriber Server". Le HSS est la base des profils des abonnés qui stocke les informations d'abonnement des utilisateurs du réseau, comme le PDN auquel ils doivent pouvoir accéder et la qualité de service dont ils doivent bénéficier.●
- L'entité PCRF signifie "Policy and Charging Rules Function" gère et contrôle le service 4 G. en gérant et en contrôlant de manière dynamique les sessions de données, ce qui permet l'émergence de nouveaux business models.
- L'entité PDN (Packet Data Network) est le réseau auquel l'UE souhaite se connecter
- L'entité PDN (Packet Data Network) est le réseau auquel l'UE souhaite se connecter

▪ **La partie IMS (IP Multimedia Sub-system) :**

L'IMS représente une architecture basée sur des nouveaux concepts et technologie qui permet de prendre en charge des sessions applicatives en temps réel telles que la visioconférence ainsi que des sessions non temps réel. En outre, l'IMS est également appelé NGN (Next Generation Network). Le cadre architectural de l'IMS se compose de 3 couches différentes services/application, contrôle et transport.

Ces 3 couches fournissent des fonctions de gestion des signaux et du trafic pour les applications multimédia [11].

I.1.1.8 Le réseau LTE-A :

LTE Advanced, ou Long-Term Evolution Advanced, est une version améliorée du standard LTE. Elle utilise des techniques plus sophistiquées pour combiner plusieurs bandes de fréquence, ce qui permet au réseau d'être beaucoup plus performant. Résultat : des débits de données qui peuvent atteindre jusqu'à 2 Gbit/s, soit environ 14 fois plus rapide que le LTE classique. C'est une grosse avancée pour profiter de connexions plus rapides et plus stables.

Les réseaux mobiles de cinquième génération (5G) :

La 5G est une technologie de communication mobile postérieure à 2020 qui offrira une connectivité élevée grâce à la technologie des commutateurs et des routeurs [12]. Ces dernières années, de nombreuses fondations de recherche et des partenaires industriels ont étudié le concept d'un réseau mobile de 5ème génération (5G) améliorant la capacité, la latence et la mobilité. [13]

Les principaux objectifs de la 5G seront d'améliorer la capacité des réseaux avec une meilleure couverture à moindre coût afin de satisfaire les besoins croissants des utilisateurs en matière de débits de données plus rapides et plus élevés. Les différents groupes de recherche travaillant sur les technologies futuristes de la 5G s'accordent sur un débit de pointe de 10 Gb/s pour les utilisateurs statiques, de 1 Gb/s pour les Utilisateurs mobiles et d'au moins 100 Mb/s dans les zones urbaines.

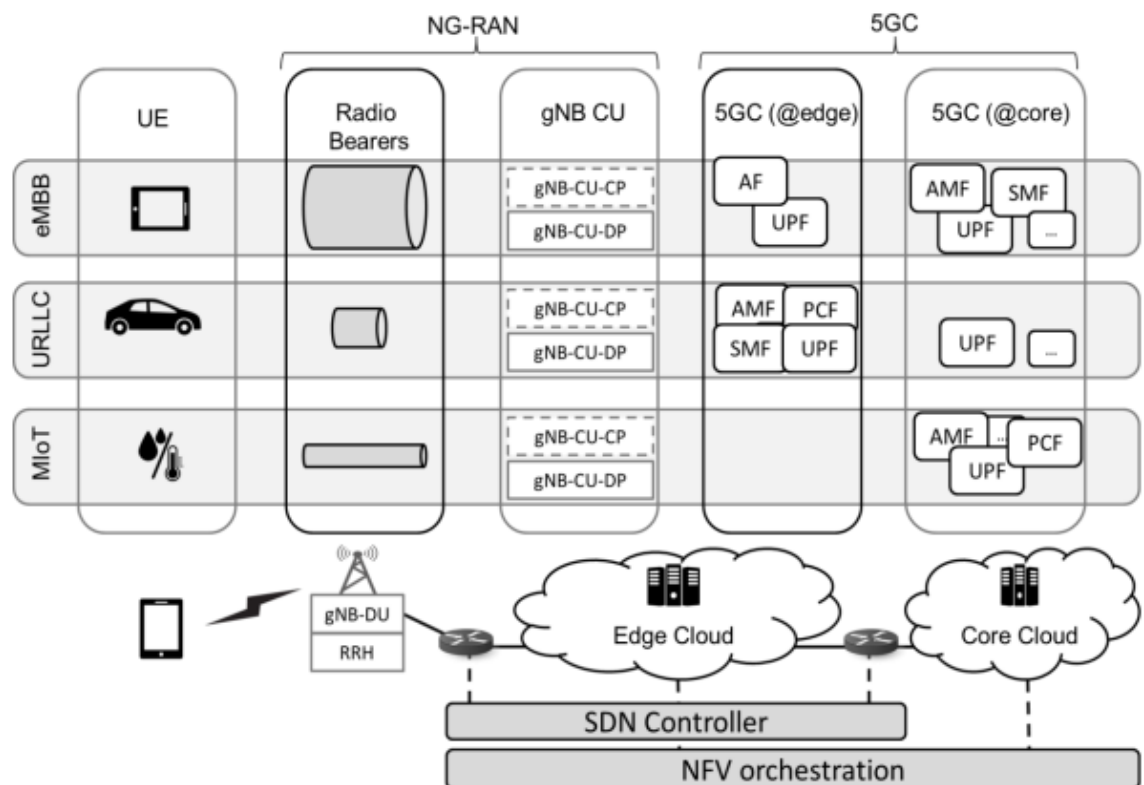


Figure I-5 : Architecture de la cinquième génération (5G).

Le NG-RAN (Next Generation Radio Access Network) est l'infrastructure radio du réseau 5G, qui se compose de stations de base 5G appelées gNB (next-generation Node B). Les gNB sont reliés au 5GC (5G Core Network) par des interfaces logiques. De la même manière que pour le LTE, les gNB peuvent être interconnectés entre eux via l'interface Xn, ce qui permet d'améliorer la mobilité des utilisateurs (comme le transfert de session) et la gestion du réseau (comme la gestion de la largeur de bande).

La fonctionnalité d'un gNB est parfois distribuée. Dans ce cas, l'architecture est constituée d'une unité centrale (gNB-CU) qui contrôle une ou plusieurs unités distribuées (gNB-DU) par l'intermédiaire de l'interface F1. Une unité distribuée est connectée à une tête radio distante (RRH), c'est-à-dire à l'émetteur-récepteur radio proprement dit. L'unité centrale est l'unité centrale est à nouveau divisée en deux parties, l'une pour les fonctions du plan de contrôle (gNB-CU-CP) et l'autre pour les fonctions du plan de donnée

(gNB-CU-CP). Fonctions du plan de donnée (gNB-CU-UP), conformément à l'approche de séparation des plans de contrôle et d'utilisation (CUPS) / SDN [14].

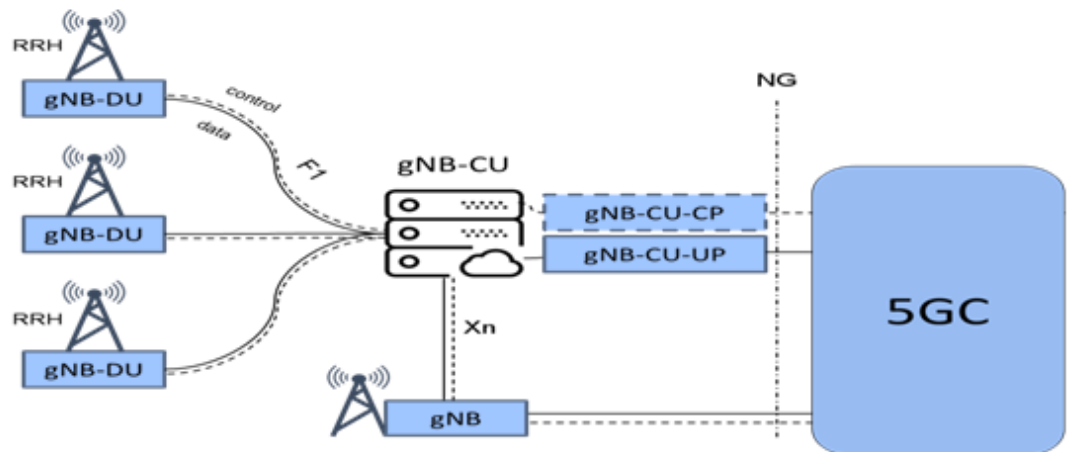


Figure I-6 : Architecture du NG-RAN

I.1.1.9 Réseau cœur 5G (5GC) :

Les principales fonctions de réseau (NFS et leurs capacités, telles qu'elles sont définies aujourd'hui dans le processus de normalisation, sont les suivantes [15] [16] :

Fonction Gestion des Accès et de la Mobilité (AMF) : Elle assure la terminaison de la signalisation NAS, le chiffrement NAS est la protection de l'intégrité, la gestion des enregistrements, la gestion des connexions, la gestion de la mobilité, l'authentification et l'autorisation d'accès, la gestion du contexte de sécurité. L'AMF comprend également la fonction de sélection des tranches de réseau (NSSF) et sert de point de terminaison pour les interfaces CP RAN (N2).



Figure I-7 : Architecture du 5G Cœur [6].

❖ **Fonction de gestion de session (SMF) :**

Elle gère l'établissement, la modification et la libération des sessions utilisateur, l'attribution et la gestion des adresses IP, ainsi que les fonctions DHCP. Elle assure aussi la signalisation NAS liée à la session, la notification des données descendantes (DL) et l'orientation du trafic vers l'UPF pour un acheminement correct.

❖ **Fonction du plan utilisateur (UPF) :**

Elle s'occupe du routage et du transfert des paquets, de l'inspection des paquets, du traitement de la qualité de service (QoS) et fait le lien avec le réseau de données (DN). Elle sert aussi de point d'ancrage pour la mobilité intra- et inter-RAT (Radio Access Technology).

❖ **Fonction d'exposition au réseau (NEF) :**

Elle expose les capacités du réseau et les événements à des applications externes. Elle permet aux utilisateurs externes (comme les entreprises partenaires) d'accéder et d'appliquer des politiques sur le réseau 5G. Le NEF assure également la sécurité dans les interactions avec les nœuds 5GC.

❖ **Fonction de référentiel réseau (NRF) :**

Le NRF gère la découverte des instances des fonctions de réseau, fournissant des informations sur les fonctions disponibles et leurs profils. Ce rôle n'existait pas dans la 4G et est essentiel pour la gestion dynamique des services du réseau 5G.

❖ **Fonction de contrôle des politiques (PCF) :**

Le PCF met en œuvre un cadre de gestion des politiques réseau, fournissant des règles pour la gestion de la mobilité, le découpage du réseau (network slicing), et la gestion de la qualité de service. Il est similaire à la fonction PCRF dans la 4G.

❖ **Gestion unifiée des données (UDM) :**

L'UDM stocke les profils des abonnés et gère les processus d'authentification et d'autorisation, ainsi que la gestion des abonnements. Il est responsable de la génération des clés de sécurité pour les utilisateurs.

❖ **Fonctions d'application (AF) :**

Les AF sont des serveurs d'application qui peuvent interagir avec les fonctions de réseau pour fournir des services. Les AF fiables peuvent accéder directement aux fonctions réseau, tandis que celles non fiables ou de tiers passent par le NEF.

❖ **Réseau de données (DN) :**

Le DN représente l'architecture du cœur du réseau 5G, permettant aux fonctions de réseau autorisées d'accéder à leurs services via des API HTTP. Ce modèle améliore la flexibilité, l'extensibilité, et facilite l'intégration avec des logiciels tiers, ce qui améliore la qualité de service.

I.3 Comparaison entre les 5 générations de communication :

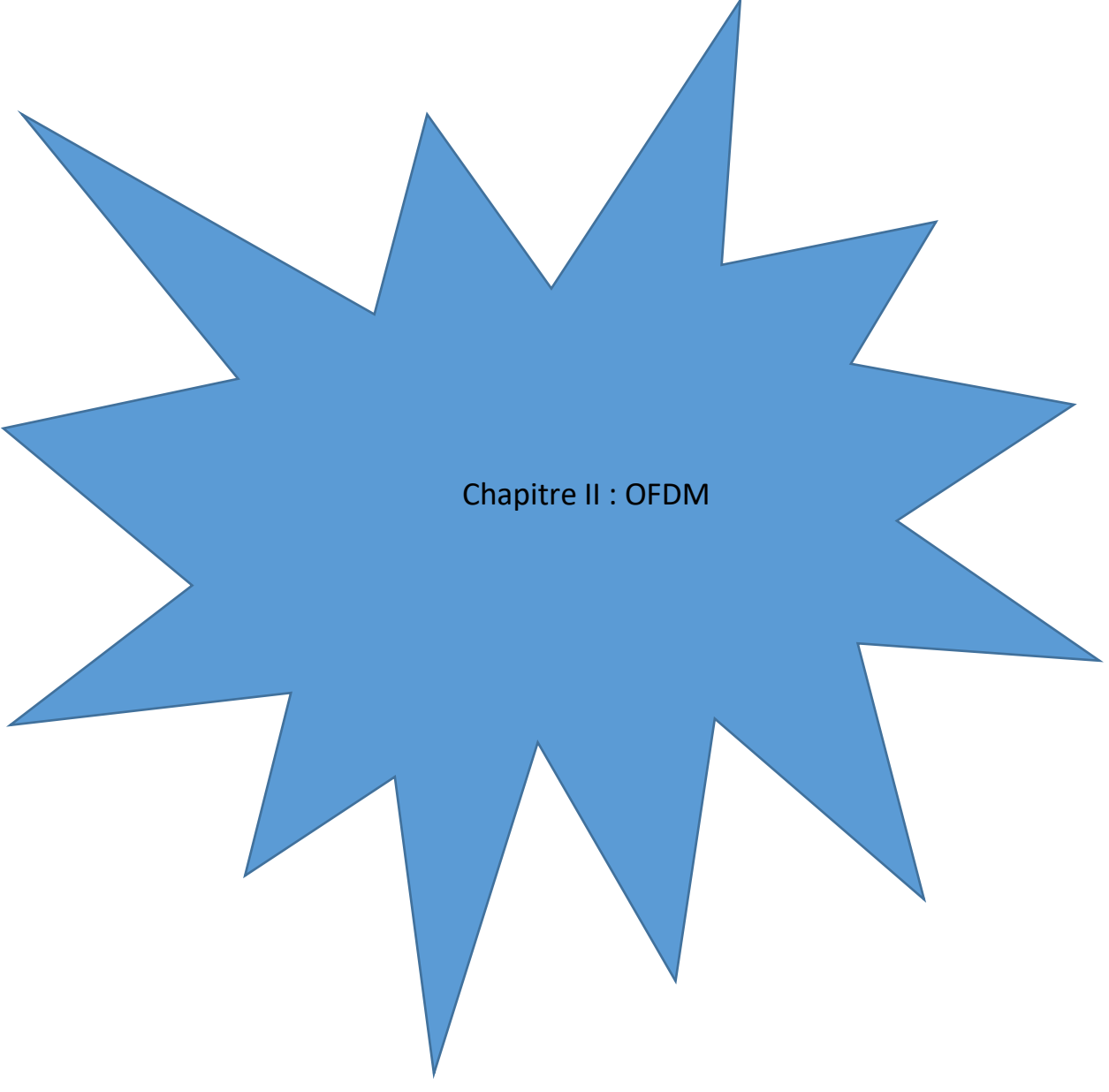
	Intitulé	Norme	Services	Année	Débit (Mbit/s)
1G	Système avancé de téléphonie mobile	AMPS	Téléphonie	1970/1980	2.4kps
	Total Access Communication système	TACS	téléphonie		
	Téléphone mobile nordique	NMT			
2G	Global System for Mobile Communication	GSM900	Digital voice, messages	1990/2004	64Kbps
	Système cellulaire numérique	DCS 1800	message		
	Service mondial de radiocommunication par paquets	GPRS			
	Enhanced Data Rate for GSM Evolution	EDGE			
3G	Universal Mobile Télécommunication système	UMTS	Intégré haute qualité audio, vidéo et de données	2004/2010	144 Kbps - 2Mbps
	Accès par paquets à haut débit (HSDPA/HSUPA)	HSPA			
	Accès par paquets à haut débit +	HSPA+			
4G	Longue Terme Evolution	LTE	Informations dynamiques	2010/2015	100 Mbps - 1Gbps
	Longue Terme Evolution Advanced	LTE_ Advanced	d'information accès, variable		

			dispositifs		
5G	Nouvelle radio	NR	Informations dynamiques d'information accès, variable dispositifs avec D'IA d'intelligence artificielle	2015/2020	10Gbps

Tableau I-2 : Comparaison entre les 5 générations

I.4 Conclusion :

Au cours de ce chapitre, nous avons examiné de manière générale l'évolution des générations de téléphones mobiles, en mettant en lumière les architectures fondamentales et les caractéristiques principales de chacune d'elles. Nous avons commencé par la première génération (1G), qui reposait sur une technologie analogique, puis avons abordé la deuxième génération (2G) avec l'introduction du GSM, un standard qui a servi de fondation aux générations suivantes. Nous avons ensuite présenté la troisième génération (3G), en détaillant ses différentes normes, telles qu'UMTS, HSPA et HSPA+, qui ont permis une amélioration significative des débits et des services multimédia. Enfin, nous avons évoqué la transition vers le LTE, qui a jeté les bases de la cinquième génération (5G), marquant ainsi un progrès majeur dans la capacité des réseaux mobiles.



Chapitre II : OFDM

II.1 Introduction :

Il est important d'avoir une compréhension fondamentale du multiplexage par répartition orthogonale de la fréquence (OFDM) car cette technologie est un élément de base pour de nombreuses normes telles que Digital Audio Broadcasting DAB, Digital vidéo Broadcasting terrestre DVB-T, Asynchrones Digital Subscriber Line ADSL comme elle est utilisée dans les différentes normes des réseaux sans fils 802.11 le 802.16 WIMAX et HiperLAN ainsi que les nouvelles générations de réseaux mobiles LTE, 4G et la 5G [17]. Dans ce chapitre, nous présentons le principe de fonctionnement de cette technique, et nous détaillons son architecture de base qui sera utilisée par la suite dans notre simulation.

II.2 Principe de fonctionnement :

Le principe de multiplexage orthogonal par répartition en fréquence OFDM est de diviser le signal numérique à transmettre en plusieurs flux de données parallèles et chaque flux est modulé sur une sous-porteuse distincte. Les sous-porteuses sont ensuite combinées pour former le signal composite OFDM, qui est transmis sur le canal de communication. L'OFDM utilise des porteuses orthogonales entre elles pour maintenir les fréquences porteuses aussi proches que possible afin de transmettre le maximum d'informations dans une bande de fréquence donnée. L'orthogonalité autorise

les signaux sur des porteuses différentes à se chevaucher sans interférer entre eux.

La figure suivante illustre les principaux concepts d'un signal OFDM et l'interrelation entre les domaines fréquentiel et temporel.

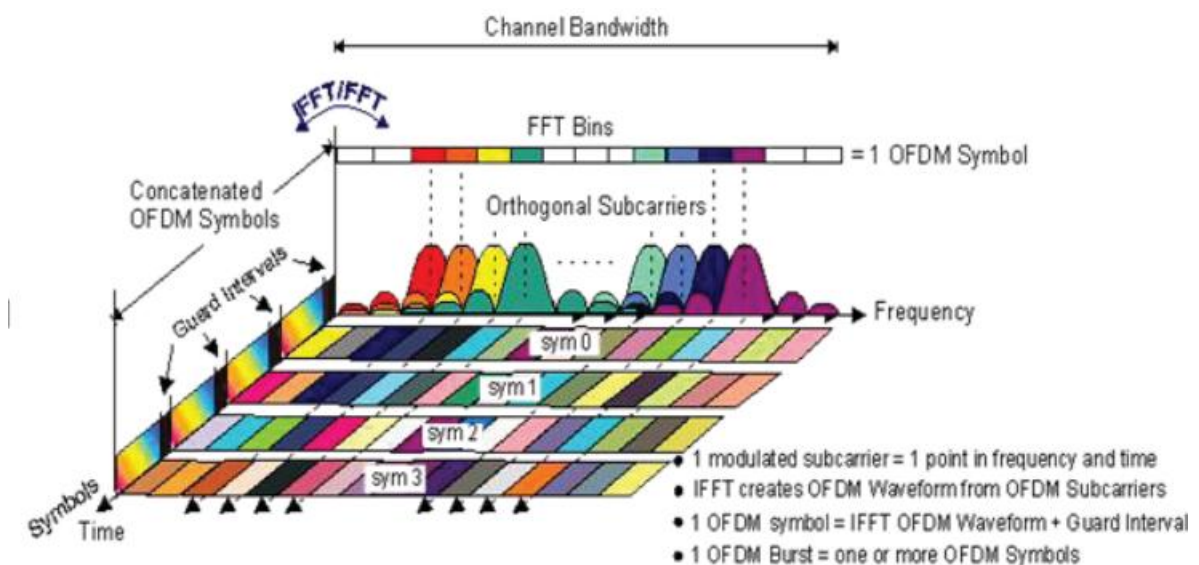


Figure II-1 : Représentations fréquentielle et temporelle du signal OFDM

L'OFDM exploite un ensemble important de sous-porteuses orthogonales rapprochées qui se transmettent parallèlement. Les sous-porteuses sont modulées en utilisant des schémas de modulation

numérique classiques tels que 16QAM QPSK, etc. à faible débit de symboles. Cependant, la combinaison de nombreuses sous-porteuses permet des débits de données similaires aux schémas de modulation à porteuse unique conventionnels dans des largeurs de bande équivalentes [18].

II.3 Chaîne de transmission :

Comme pour tout modèle de transmission, la chaîne OFDM (Orthogonal Frequency Division Multiplexing) se divise en trois grandes parties principales : l'émetteur, le canal et le récepteur. Chacune de ces parties comprend plusieurs blocs fonctionnels spécifiques qui sont responsables de l'ensemble du processus de transmission et de réception des données. Dans la prochaine sous-section, nous détaillerons le rôle de chaque bloc, en expliquant leur fonctionnement et leur interaction au sein de la chaîne OFDM.

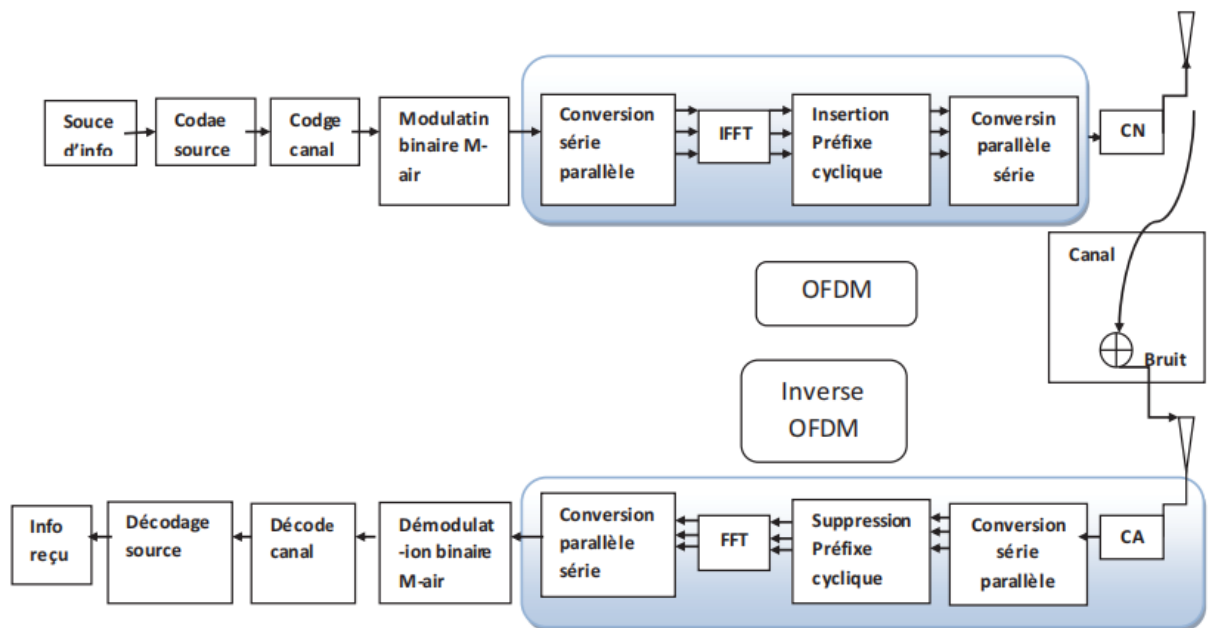


Figure II-2 : Diagramme en bloc de la chaîne de transmission OFDM

Émetteur :

Source d'information : au cours de cette étape, les capteurs jouent un rôle essentiel en transformant les informations physiques telles que les ondes sonores, lumineuses ou thermiques en signaux Électriques. Ces signaux sont ensuite convertis en une séquence de bits d'information [19].

Codage source : La fonction de cette étape est de réduire la quantité de données en éliminant les Informations redondantes en bits en les compressant.

Codage canal : Le principe du codage est d'ajouter des bits supplémentaires aux données à transmettre afin de protéger le signal transmis sur le canal contre les erreurs. Le récepteur utilise ces bits pour détecter et corriger les erreurs.

Modulation M-air : Consiste à modifier les propriétés d'un signal périodique appelé signal porteur $S(t) = A \cos(\omega_0 t + \varphi_0)$ avec un signal de modulation qui contient des données à transférer. Ceci

est

fait par la modification d'un ou plusieurs paramètres qui sont l'amplitude, la phase et la fréquence.

Le signal modulé s'écrit comme suit : [20]

$$m(t) = R \sum_k x_k(t) e^{j2\pi f_0 t} \quad (\text{II- 1})$$

Où :

$x_k(t)$: Est l'enveloppe complexe correspond au $k^{\text{ème}}$ symbole :

$$x_k(t) = \sum x_k \text{ rect}(t - kT) \quad (\text{II- 2})$$

Il existe plusieurs types de technique de modulation numérique disponible en fonction de l'exigence.

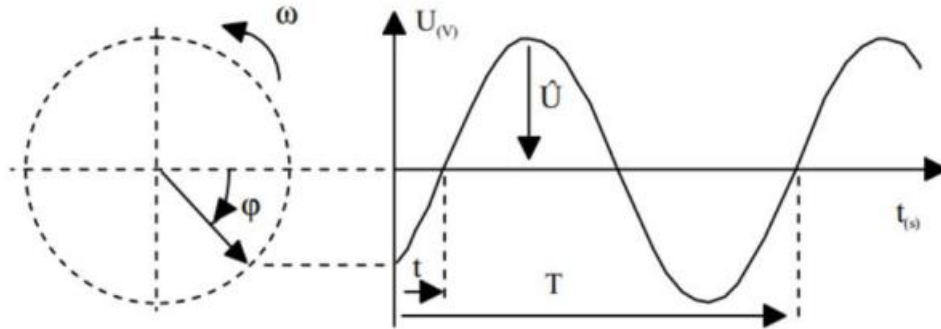


Figure II-3 représentation temporelle et vectorielle [21].

En variant $\hat{U}$ c'est une modulation par déplacement d'amplitude (**Amplitude Shift Keying (ASK)**).

En variant f (donc la pulsation ω) c'est une modulation par déplacement de fréquence (**Frequency Shift Keying (FSK)**).

En variant ϕ c'est une modulation par déplacement de phase (**Phase Shift Keying (PSK)**).

QAM (Quadrature Amplitude Modulation) est obtenue en variant l'amplitude de la porteuse et d'une onde en quadrature.

Modulation OFDM :

(Orthogonal Frequency Division Multiplexing) est une méthodologie de modulation appliquée dans les réseaux de communication sans fil pour transmettre données sur plusieurs sous porteuses orthogonales. En lieu et place d'utiliser une seule porteuse pour transmettre toutes les données, l'OFDM sépare le signal en une série de sous-porteuses, où chaque sous-porteuse transporte une partie des données. Chaque sous-porteuse est modulée indépendamment à une fréquence déterminée, puis les signaux modulés sont réunis pour créer le signal OFDM intégral.

À cette étape les données subissent à une conversion série/parallèle ensuite à une transformée de Fourier pour passer du temporelle au fréquentielle.

La IFFT est définie par :

$$X[n] = \frac{1}{N} \sum_{k=0}^{N-1} X[k] e^{\frac{j2\pi kn}{N}}, \quad 1 \leq n \leq N$$

Préfixe cyclique :

Dans les systèmes OFDM la propagation multi-trajets provoque une interférence entre symbole et pour éviter ce problème on utilise la bande de garde, le matériel ne permet pas d'espace vide car il nécessite d'envoyer des signaux en continu. Donc on utilise le préfixe cyclique qui est réalisé par l'ajoute d'un intervalle entre chaque symbole OFDM utile à la sortie de l'IFFT à l'émission, et dans la réception le préfixe cyclique est supprimé avant d'effectuer la FFT. Le préfixe cyclique est une copie de l'extrémité du symbole OFDM situé au début de ce symbole.

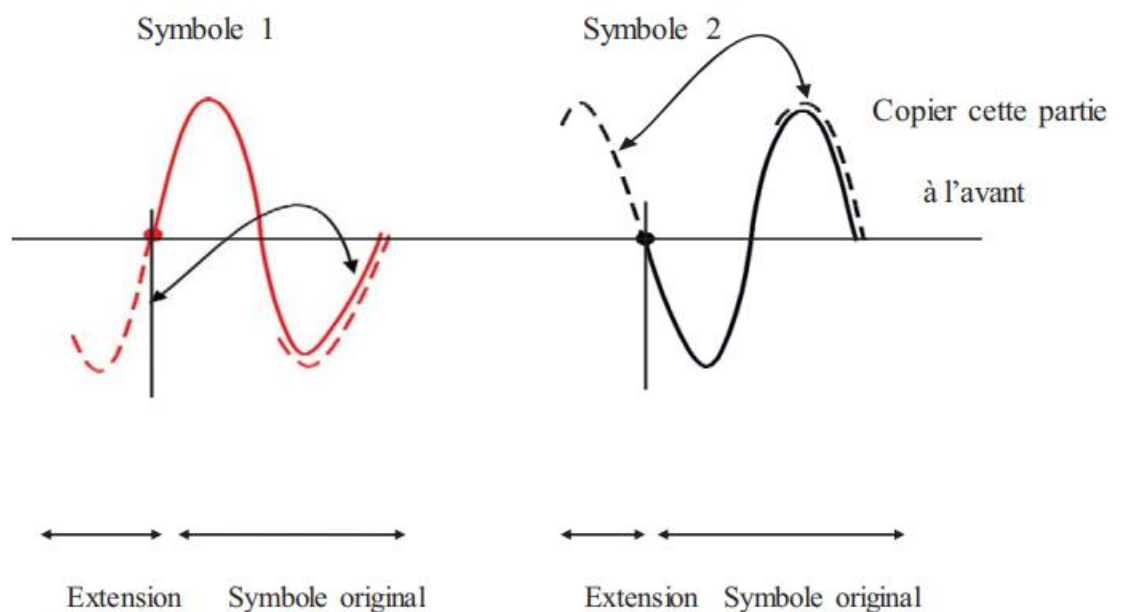


Figure II-4 : Préfixe cyclique

Canal de transmission :

Un canal de transmission est un support physique ou une voie de communication qui permet de transmettre les informations entre l'émetteur et le récepteur. Il peut être un câble coaxial, une fibre optique, une liaison hertzienne ou bien une liaison sans fils tel qu'il est dans les réseaux cellulaires. Dans cette section, nous aborderons les canaux de communication sans fil en exposons les diverses caractéristiques de ces canaux.

Le type de canal de propagation traversé par un signal peut avoir un impact significatif sur la force, la qualité et la fiabilité du signal. Voici quelques modèles [22] :

- ❖ **Canal en visibilité directe (LOS)** : Dans ce cas, le signal se propage en ligne droite entre l'émetteur et le récepteur, sans aucune obstruction. Cela permet une transmission claire et optimale, car il n'y a pas de perturbations dues à des objets physiques sur le trajet du signal.
- ❖ **Canal sans visibilité directe (NLOS)** : Ici, il y a des obstacles physiques (comme des bâtiments, des montagnes ou d'autres objets) entre l'émetteur et le récepteur, empêchant une ligne de visée directe. Le signal peut encore atteindre le récepteur, mais il sera réfléchi ou réfracté par ces obstacles, ce qui peut causer des interférences ou une dégradation du signal.
- ❖ **Canal à trajets multiples** : Ce phénomène se produit lorsqu'un signal atteint le récepteur par plusieurs chemins différents. Ces chemins peuvent être causés par des réflexions, des diffractions ou des diffusions. Ces trajets multiples peuvent interférer les uns avec les autres, provoquant des effets tels que le "fading" (affaiblissement) du signal, ce qui rend la réception plus complexe.
- ❖ **Fading Channel** : Ce type de propagation est lié aux changements dans l'environnement, comme le mouvement des objets (voitures, personnes, etc.) ou des conditions atmosphériques (pluie, neige, etc.). Cela cause une variation de l'intensité du signal reçu, parfois de manière rapide et imprévisible, et peut affecter la qualité de la communication.

Dans notre étude on s'intéresse au canal à trajet multiples. Les voies de propagation sont à trajets multiples en raison d'obstacles autour de l'émetteur et du récepteur (figure II.5). Dans cette situation, le récepteur reçoit plusieurs copies du signal transmis qui ont emprunté différents chemins, présentant des atténuations, des déphasages et des retards variés (en raison de la longueur des chemins).

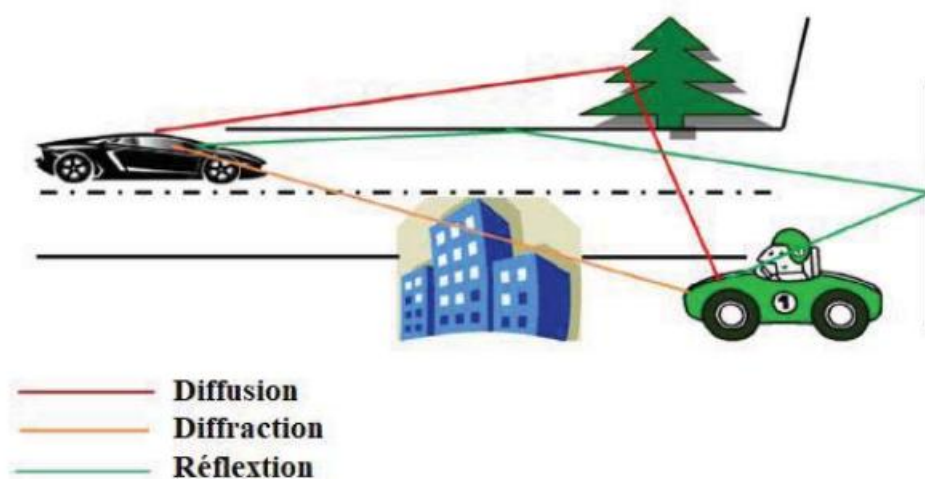


Figure II-4 : propagation par trajets multiples

La propagation multi-trajets dans les communications sans fil repose sur trois mécanismes << comme une structure artificielle, et est réfléchi. Le deuxième c'est la diffraction, qui se manifeste

lorsque le trajet est obstrué sur une grande partie de sa circonférence par rapport à la longueur d'onde du signal d'entrée. La troisième c'est la diffusion, qui se manifeste lorsque le signal atteint le bord d'une structure artificielle, entraînant ainsi une dispersion de l'énergie du signal dans diverses directions [23].

II.1.1.1 Canal de propagation par trajets multiples :

Il existe de nombreux modèles mathématiques de canaux d'évanouissement décrits dans la littérature, qui décrivent les caractéristiques statistiques. Parmi les distributions les plus couramment utilisées figurent la distribution de Rayleigh et la distribution de Rice.

▪ Distribution de Rayleigh et Rayleigh fading :

Les retards correspondant aux variés chemins de signaux dans un canal à évanouissement par trajets multiples sont imprévisibles et se peuvent décrire que statistiquement.

Si il y a beaucoup de chemins, le théorème central limite peut être utilisé pour modéliser la réponse impulsionnelle aléatoire avec temps par le canal en tant que processus aléatoire

gaussien à valeur complexe. Lorsque la réponse impulsionnelle est modélisée comme un processus gaussien à valeur complexe de moyenne nulle, on dit que le canal est un canal à évanouissement de Rayleigh[24].

Ici, le modèle d'évanouissement de Rayleigh est supposé n'avoir que deux composantes multivoie $X(t)$ et $Y(t)$. En additionnant simplement les deux variables aléatoires gaussiennes et en prenant la racine carrée (enveloppe), on obtient un processus distribué de Rayleigh à une seule prise. La phase d'une telle variable aléatoire suit une distribution uniforme.

Considérons deux variables aléatoires gaussiennes de moyenne nulle et de même variance $X \sim (0, \sigma^2)$ et $Y \sim (0, \sigma^2)$.

Définissons une variable aléatoire gaussienne complexe :

$$Z = X + jY \quad (II- 4)$$

L'enveloppe de la variable aléatoire complexe est donné par :

$$R = \sqrt{X^2 + Y^2} \quad (II- 5)$$

et la phase est donné par :

$$\varphi = \tan^{-1}\left(\frac{Y}{X}\right) \quad (II- 6)$$

L'enveloppe suit la distribution de Rayleigh et la phase sera uniformément distribuée. La fonction de densité de probabilité (distribution de Rayleigh) de la réponse en amplitude mentionnée ci-dessus est donnée par :

$$P_R(r) = \frac{r}{\sigma^2} e^{-\left(\frac{r^2}{2\sigma^2}\right)}, \quad r \geq 0 \quad (II- 7)$$

- **.Distribution de Rice et Rice fading :**

Le modèle de l'évanouissement de Rician est similaire à celui de l'évanouissement de Rayleigh, sauf que dans l'évanouissement de Rician, une forte composante dominante est présente. Cette composante dominante peut par exemple être line-of-sight wave. Cette composante est modélisée à l'aide de deux variables aléatoires gaussiennes, l'une à moyenne nulle et l'autre à moyenne non nulle.

Puisque les deux variables X et Y ont des "moyennes" différentes, un paramètre de non centralité (indiquant la moyenne non centrale) est défini :

$$s = \sqrt{m_1^2 + m_2^2} \quad (\text{II- 8})$$

Le paramètre de non centralité (le déséquilibre des moyennes) est causé par la présence d'un chemin dominant dans un environnement d'évanouissement de Rice. Pour cette raison, le facteur K de Rice - représentant le rapport entre la puissance de la Line-Of-Sight (LOS) et la puissance de Non-Line-Of Sight (NLOS) est défini dans un tel scénario.

$$K = \frac{\text{Power of LOS component}}{\text{Power of NLOS component}} \quad (\text{II- 9})$$

Statistiquement, cela peut être représenté comme la puissance dans l'enveloppe délavée qui a été produite par les moyens de X et Y.

$$K = \frac{m_1^2 + m_2^2}{2\sigma^2} = \frac{s^2}{2\sigma^2} \quad (\text{II- 10})$$

L'enveloppe suit la distribution de Rician, dont la PDF est donnée par :

$$p(r) = \frac{r}{\sigma^2} e^{\left(-\frac{r^2+s^2}{2\sigma^2}\right)} I_0\left(\frac{sr}{\sigma^2}\right), \quad r \text{ et } s \geq 0 \quad (\text{II- 11})$$

Où $I_0(x)$ est la fonction de Bessel modifiée de première espèce d'ordre 0.

II.1.1.2 Canal à Bruit blanc additif gaussien (AWGN)

Le canal Le canal BBAG est considéré comme un canal de radio mobile classique et simple. Il est caractérisé par la présence d'un bruit blanc additif qui résulte de l'addition d'un signal émis et d'un bruit blanc. Ce type de canal est souvent utilisé dans les systèmes mobiles statiques où le signal émis n'est pas altéré, à l'exception du bruit. Le bruit additif est modélisé à l'aide d'un procédé aléatoire, blanc, gaussien et centré, et est indépendant du signal. La densité spectrale bilatérale de puissance du bruit est constante [26] [27], ce qui permet de considérer que le canal est sans distorsion.

$$y(t) = A X(t) + b(t) \quad (\text{II- 12})$$

Le signal modulé est représenté par $x(t)$ et l'affaiblissement général sur le trajet est noté A, supposé constant dans le temps. Le bruit présent dans le système est quant à lui modélisé par un processus aléatoire

gaussien $b(t)$, également appelé bruit blanc Gaussien AWGN. Ce bruit est caractérisé par une moyenne nulle, une variance et une densité spectrale de puissance bilatérale $\frac{N_0}{2}$, La densité de probabilité donnée par :

$$p_Y\left(\frac{Y}{X}\right) = \frac{1}{\sqrt{2\pi\sigma_b^2}} e^{-\frac{(y-x)^2}{2\sigma_b^2}} \quad (\text{II- 13})$$

Et la puissance de bruit moyenne est écrite comme suit :

$$p_n = \frac{E[x^2(t)] + E[y^2(t)]}{2} = \frac{\sigma_b^2 + \sigma_b^2}{2} = \sigma_b^2 \quad (\text{II- 14})$$

II.1.1.3 Canal de propagation à bruit impulsif

Le bruit impulsif, par sa nature non stationnaire et non gaussienne, présente un défi majeur pour les systèmes de communication modernes. Il peut être causé par des sources non naturelles de rayonnement (par exemple, des équipements électroniques ou industriels), ce qui entraîne des perturbations plus difficiles à gérer par rapport au bruit thermique classique. La gestion de ce type de bruit nécessite des techniques spécifiques, comme le filtrage avancé et la correction d'erreurs, pour préserver la qualité de la transmission du signal.

Le bruit impulsif se manifeste dans divers phénomènes de propagation tels que le bruit reçu par un sonar lors de la propagation sous-marine, le bruit rencontré lors de la réception d'un signal radar en raison des échos multiples, ainsi que les perturbations électromagnétiques causées par les orages. En outre, le bruit impulsif peut également être généré par des manœuvres aléatoires d'appareils électriques tels que les redresseurs commandés, les moteurs électriques et les appareils de commutation, qui sont connus pour créer des bruits impulsifs dans les environnements domestiques ou en intérieur [28].

Dans la littérature récente, plusieurs modèles ont été employés pour simuler des phénomènes non gaussiens, qui étaient initialement étudiés par Middleton. Au début, le modèle de Bernoulli Gaussien a été utilisé pour modéliser la séquence de bruit impulsif, où l'amplitude est décrite par une distribution gaussienne et le taux d'occurrence des impulsions est modélisé par une distribution de Bernoulli [24]. Cependant, les chercheurs ont remis en question la pertinence de cette modélisation du bruit avec une distribution gaussienne en raison de sa grande variabilité. Par conséquent, il est plus approprié de modéliser ce bruit avec une distribution qui n'admet pas nécessairement une variance finie plutôt qu'avec une distribution gaussienne. Les distributions alpha-stables sont apparues comme une alternative aux distributions gaussiennes à variance infinie, offrant ainsi une modélisation plus précise des environnements d'interférence impulsive. Ils appartiennent à une classe riche en probabilités et en lois de Gauss de Cauchy et de Lévy.

- **Middleton classe A :**

Dans cette section, nous abordons en détail le modèle de bruit de classe A, qui est largement utilisé dans la modélisation du rayonnement impulsif et a fait l'objet de nombreuses études dans la littérature [29].

Ce modèle définit la fonction de densité de probabilité (PDF) d'un échantillon de bruit de la manière suivante [30] :

$$F_m(n_k) = \sum_{m=0}^{\infty} p_m N(n_k; 0; \sigma_m^2) \quad (\text{II- 15})$$

Où $N(n_k; 0; \sigma_m^2)$ est le PDF gaussien avec moyenne $\mu = 0$ et écart de σ_m^2 .

$$p_m = \frac{A^m e^{-A}}{m!} \quad (\text{II- 16})$$

m : est le nombre d'interférences actives et A représenté la densité des impulsions dans une période d'observation donnée par la relation suivante :

$$A = \frac{(\eta\Gamma)}{T_0} \quad (\text{II- 17})$$

Où T_0 est le temps unitaire et il est égal à un, η représente le nombre moyen d'impulsions par seconde et Γ est la moyenne durée de chaque impulsion, où toutes les impulsions sont prises ont la même durée [31].

Le paramètre A d'écrit le bruit comme suit : Quand A diminué le bruit devient plus impulsif, si non le bruit tend vers AWGN.

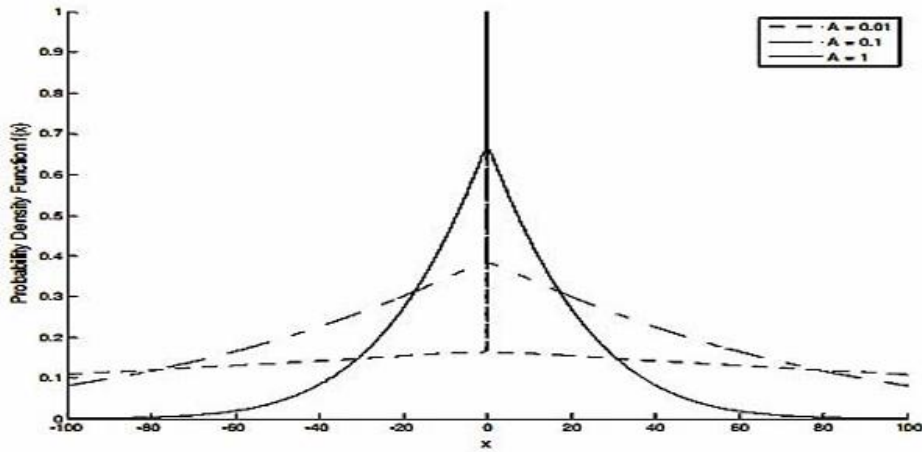


Figure II-5 : Bruit pour différentes valeurs de A et $\Gamma=0.001$

En revanche, σ_m peut-être écrit comme :

$$\sigma_m = \sigma_2 \frac{\frac{m}{A} + \Gamma}{1 + \Gamma} \quad (\text{II- 18})$$

Avec :

$$\sigma = \sigma_G^2 + \sigma_C^2 \quad (\text{II- 19})$$

est la puissance sonore totale a deux composantes : une gaussienne variance σ_G^2 et les autres variations de bruit impulsif σ_i^2 et :

$$\Gamma = \frac{\sigma^2 G}{\sigma_i^2} \quad (\text{II- 20})$$

est le rapport de puissance de bruit gaussien-impulsif. Il peut être observé que pour les valeurs basses de, la composante impulsive et pour les valeurs supérieures, le composant AWGN prédomine [32]

▪ **Bernoulli-Gaussien :**

Après avoir présenté le modèle de bruit de classe A de Middleton au paragraphe précédent, nous allons maintenant aborder un autre modèle couramment utilisé appelé modèle de bruit Bernoulli Gaussien, qui est mentionné dans [33], [26]. Ce modèle représente la somme de deux fichiers PDF gaussiens pondérés par une distribution Bernoulli. Dans le cadre d'un modèle Bernoulli-Gaussien pour un processus de bruit impulsif, le temps d'apparition aléatoire des impulsions est modélisé par un processus binaire de Bernoulli $b(m)$, tandis que l'amplitude des impulsions est modélisée par un processus gaussien $n(m)$. Un processus de Bernoulli $b(m)$ est un processus binaire ce qui amène la valeur 1 avec une probabilité p et la valeur 0 avec une probabilité de $1-p$. Ce modèle de bruit est décrit par le PDF suivant :

$$FBG(n_k) = (1 - P)N(n_k; 0; \sigma_G^2) + p N(n_k; 0; \sigma_G^2 + \sigma_i^2) \quad (\text{II- 21})$$

bien que les deux modèles partagent une structure similaire de perturbation des données, la manière dont le bruit est ajouté et géré diffère entre le modèle Bernoulli-gaussien et le modèle de bruit de classe A de Middleton.

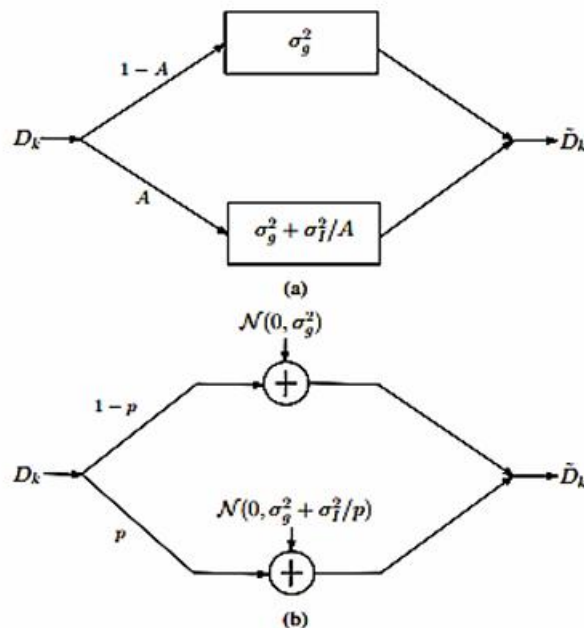


Figure II-6 : (a) deux modèles de bruit de classe d'état et (b) modèle de bruit Bernoulli-gaussien

▪ **Alpha α distribution stable :**

n plus des deux modèles de bruit impulsif déjà discutés — le modèle de bruit de classe A de Middleton et le bruit Bernoulli-Gaussien — un autre modèle de bruit de plus en plus utilisé dans la littérature est la distribution symétrique α stable ($S\alpha S$) [34][35]. Nous considérons une variable X qui suit une loi α -stable de paramètres : α, β, γ et δ ; $X \sim S\alpha(\beta, \gamma, \delta)$ si et seulement si sa caractéristique est de la forme :

$$\Psi_{\alpha(t)} = \exp\{-\gamma^\alpha |t|^\alpha [1 + i\beta \text{sign}(t) \omega(t, \alpha)] + i\delta(t)\} \quad (\text{II- 22})$$

Avec :

$$\omega(t, \alpha) = \begin{cases} -\tan\left(\frac{\alpha\pi}{2}\right) & \text{si } \alpha \neq 1 \\ \frac{2}{\pi} \ln|t| & \text{si } \alpha = 1 \end{cases} \quad (\text{II- 23})$$

et :

$$\text{sign}(t) = \begin{cases} 1 & t > 0 \\ 0 & t = 0 \\ -1 & t < 0 \end{cases} \quad (\text{II- 24})$$

Une distribution stable est définie par quatre paramètres : l'exposant caractéristique (α), qui contrôle la lourdeur de la queue et peut prendre n'importe quelle valeur dans l'intervalle $\alpha \in]0; 2]$, le paramètre d'échelle $\gamma > 0$, le paramètre d'emplacement (δ), et le paramètre de symétrie (β), qui ne peut prendre des valeurs que dans certain intervalle $\beta \in [-1; 1]$.

Figure II.9 montre les PDF de la α - modèles de distribution stables pour différentes valeurs de α tandis que les autres paramètres sont maintenus fixes ($\beta = 0$; $\gamma = 1$ et $\delta = 0$)

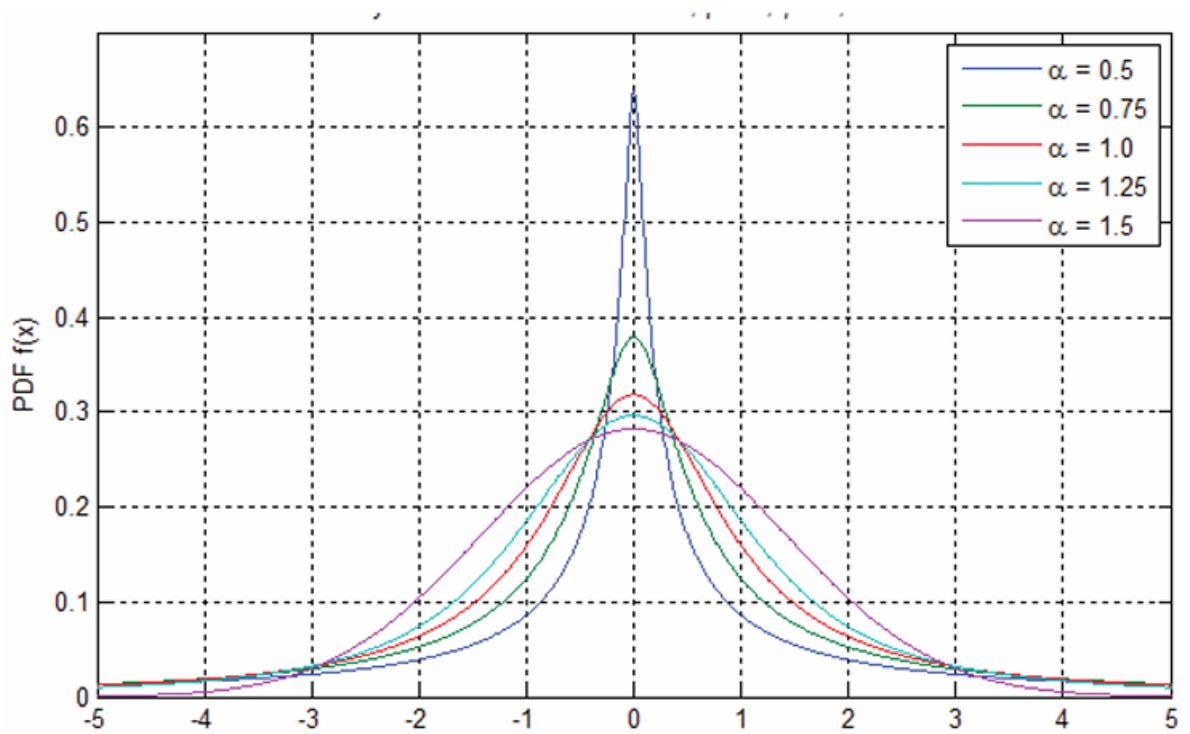


Figure II.7 : distributions de différentes valeurs de α tandis que $\beta = 0$; $\gamma = 1$ et $\delta = 0$

Le bruit impulsif ne possède pas de formule explicite pour sa fonction de densité de probabilité (PDF), sauf dans certains cas particuliers de la sous-classe de distribution symétrique α -stable ($S\alpha S$) caractérisée par $\beta = 0$. [36]. Ces cas particuliers sont résumés dans le tableau ci-dessous.

Distribution	α	Equation
Gaussien	2	$F(X) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{(x-\delta)^2}{2\sigma^2}\right)$
Cauchy	1	$F(X) = \frac{\gamma}{\pi} [\gamma^2 + (x-\delta)^2]$
Levy	0.5	$F(X)$ $= \sqrt{\frac{\gamma}{2\pi}} \frac{1}{(x-\delta)^{\frac{3}{2}}} \exp\left(-\frac{\gamma}{2(x-\delta)}\right) \delta < x < \infty$

Tableau II-1 : les équations des cas particuliers de α - stable

D'autre part il est possible d'approcher par la transformation inverse de la fonction caractéristique le PFD d'une loi stable en dehors d'une intégrale écrite comme suit :

$$F_X(X) = \frac{1}{2\pi} \int_{-\infty}^{+\infty} e^{-itx} \Psi_{\alpha}(t) dt \quad (\text{II- 25})$$

Récepteur :

La fonction de la partie réceptrice consiste à récupérer le message qui a été transmis par l'émetteur à travers le canal. La chaîne de réception est composée de différents blocs qui assurent des fonctions inverses à celles effectuées lors de l'émission. Le but est de traiter le signal de manière à ce que le message soit correctement reçu et compris par le destinataire

- ❖ **Antenne de réception** : L'antenne capte les ondes électromagnétiques transmises à travers l'air. Elle joue un rôle fondamental en transformant ces ondes en un signal électrique que le récepteur peut traiter.
- ❖ **Suppression du préfixe cyclique ou remplissage de zéros** :
 - Le préfixe cyclique (dans les systèmes OFDM) est une séquence de données copiées qui est ajoutée au début de chaque symbole pour combattre l'effet des interférences inter-symboles dues aux réflexions. Lors de la réception, il doit être supprimé avant que la démodulation ne puisse avoir lieu.
 - Alternativement, le remplissage de zéros peut être utilisé dans certains systèmes pour éviter des interférences supplémentaires.
- ❖ **Démodulation OFDM** : La modulation OFDM permet de diviser les données en plusieurs sous-porteuses. La démodulation est réalisée via une transformée de Fourier rapide (FFT), qui permet de récupérer les données originales en transformant le signal reçu dans le domaine fréquentiel. Cela revient à effectuer l'opération inverse de la modulation effectuée lors de l'émission du signal.
- ❖ **Démodulation M-aire/binaire** : Cette étape consiste à convertir les symboles reçus (qui peuvent être sous forme de modulations comme QAM, PSK, etc.) en paquets de bits. En fonction du type de modulation, cela permet de reconstruire les bits d'origine.
- ❖ **Décodage canal et de décodage source** :
 - **Décodage canal** : Cette étape utilise des techniques de correction d'erreurs, telles que les codes de correction d'erreurs (par exemple, les codes Reed-Solomon ou Turbo Codes), pour détecter et corriger les erreurs qui peuvent avoir été introduites pendant la transmission à travers le canal.
 - **Décodage source** : Cette phase vise à restaurer les données originales en supprimant les redondances qui ont été ajoutées pendant l'encodage source à l'émission. Cela peut également inclure une décompression des données si elles ont été compressées avant l'envoi.

Conclusion :

Ce chapitre a exploré la modulation OFDM, en mettant en lumière son fonctionnement et les éléments essentiels de son processus de traitement. Nous avons aussi étudié les divers types de canaux de transmission répertoriés dans la littérature, ainsi que l'introduction du bruit impulsif. Cette structure servira de base pour la suite de notre travail



Chapitre III : Apprentissage en profondeur

III.1 Introduction :

L'apprentissage profond, aussi appelé Deep Learning, c'est une façon avancée d'entraîner des ordinateurs à apprendre en utilisant des réseaux de neurones artificiels. En gros, ces réseaux ont plusieurs couches qui traitent l'information, ce qui leur permet d'analyser des données de manière très fine et de repérer des détails à différents niveaux de complexité. Cela leur donne une capacité d'apprentissage impressionnante, facilitant la compréhension et la modélisation d'informations complexes.

L'apprentissage automatique est un sous-domaine de l'intelligence artificielle (IA) qui consiste à rendre les machines aussi intelligentes que le cerveau humain. Ce chapitre fournit une description de l'apprentissage profond et les réseaux de neurones [37].

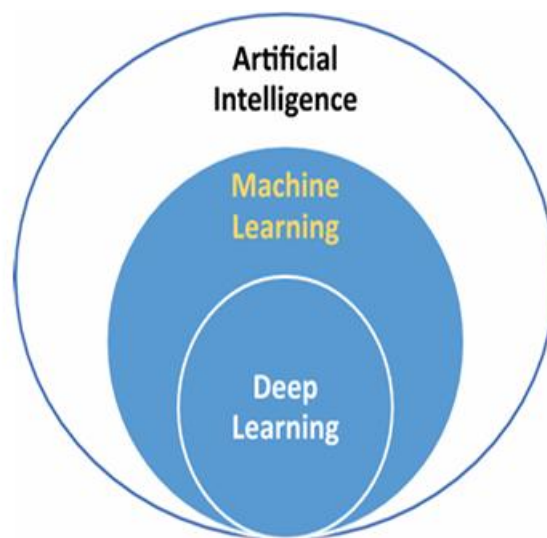


Figure III-1 : Diagramme de Venn

III.2 Les réseaux neuronaux :

Les réseaux de neurones artificiels sont essentiels dans l'apprentissage automatique et jouent un rôle central dans le Deep Learning. Inspirés du fonctionnement du cerveau humain, ils imitent la façon dont les neurones biologiques communiquent pour traiter et analyser les données de manière intelligente. Avec leur architecture composée de plusieurs couches, ils peuvent extraire des caractéristiques complexes et apprendre à résoudre différents types de problèmes plus facilement.

Présentation sommaire du neurone biologique :

Le système nerveux est composé de milliards de cellules :

C'est un réseau de neurones biologiques. En effet, les neurones ne sont pas indépendants les uns des autres, ils établissent entre eux des liaisons et forment des réseaux plus ou moins complexes.

Le neurone biologique est composé de trois parties principales :

- ❖ **Le corps cellulaire** : est composé du centre de contrôle traitant les informations reçues par les dendrites.
- ❖ **Les dendrites** : sont les principaux fils conducteurs par lesquels transitent l'information venue de l'extérieur.
- ❖ **L'axone** : est le fil conducteur qui conduit le signal de sortie du corps cellulaire vers d'autres neurones.

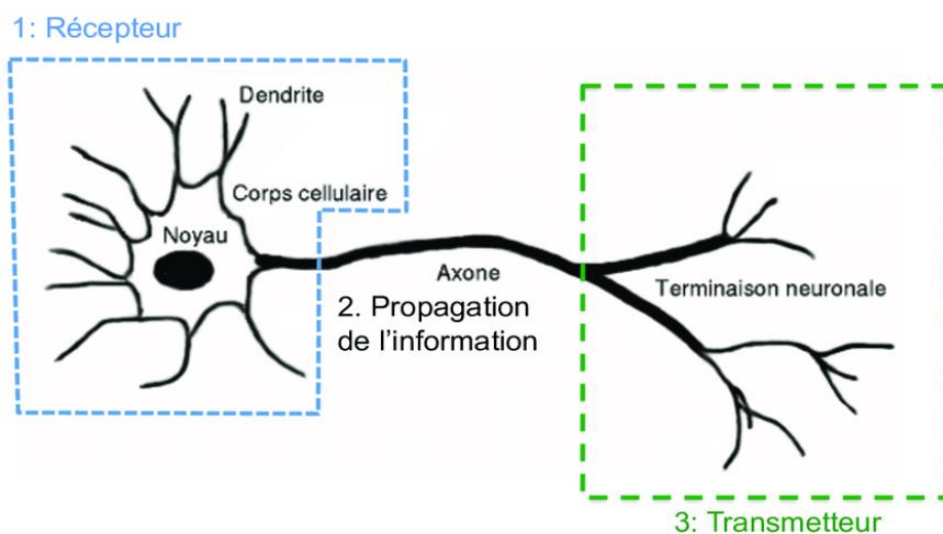


Figure III-2 : Schéma d'un neurone biologique [38]

Quant aux synapses, elles font effet de liaison et de pondération entre neurones et permettent donc aux neurones de communiquer entre eux.

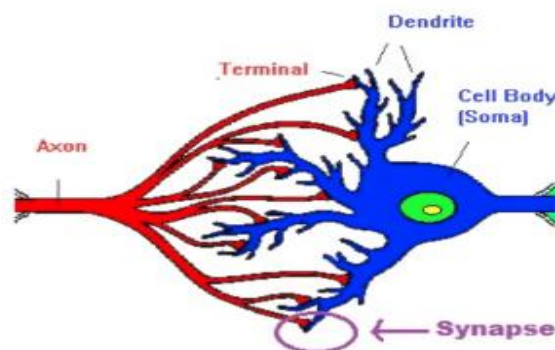


Figure III-3 : Neurone biologique

Les réseaux neuronaux artificiels :

Les réseaux neuronaux artificiels sont un peu comme des cerveaux numériques. Ce sont des systèmes qui apprennent en s'inspirant du fonctionnement des neurones dans notre cerveau. Imaginez des couches de petits nœuds, qu'on appelle neurones artificiels, qui communiquent entre eux, traitant toutes sortes

d'informations grâce à des connexions qui ont des poids, comme des réglages. En s'entraînant avec une méthode appelée rétropropagation, ces réseaux vont doucement ajuster leurs paramètres pour devenir de plus en plus précis à réaliser différentes tâches. On les retrouve dans plein de domaines, comme la reconnaissance d'images, le traitement du langage ou encore la détection de formes. Ces modèles sont capables de résoudre des problèmes difficiles et d'extraire directement les informations importantes à partir des données qu'on leur donne. Parmi leurs différentes architectures, on peut citer les CNN (réseaux de neurones convolutionnels) pour analyser des images, les RNN (réseaux récurrents) pour traiter des séquences, et les modèles de type transformeurs pour comprendre le langage. Grâce à ces avancées, ils ont changé la donne dans beaucoup de domaines, et ils continuent d'améliorer la façon dont l'intelligence artificielle progresse.

La structure d'un réseau neuronal artificiel se compose généralement d'une entrée (input layer), d'une sortie (output layer) et d'une ou plusieurs couches cachées. Chaque couche d'un réseau de neurones artificiels est constituée de plusieurs neurones. Chaque neurone est relié à tous les autres neurones des couches adjacentes. Autrement dit, il prend les sorties des couches précédentes comme entrées et les combine linéairement pour produire le stimulus qui est envoyé à la fonction d'activation [39].

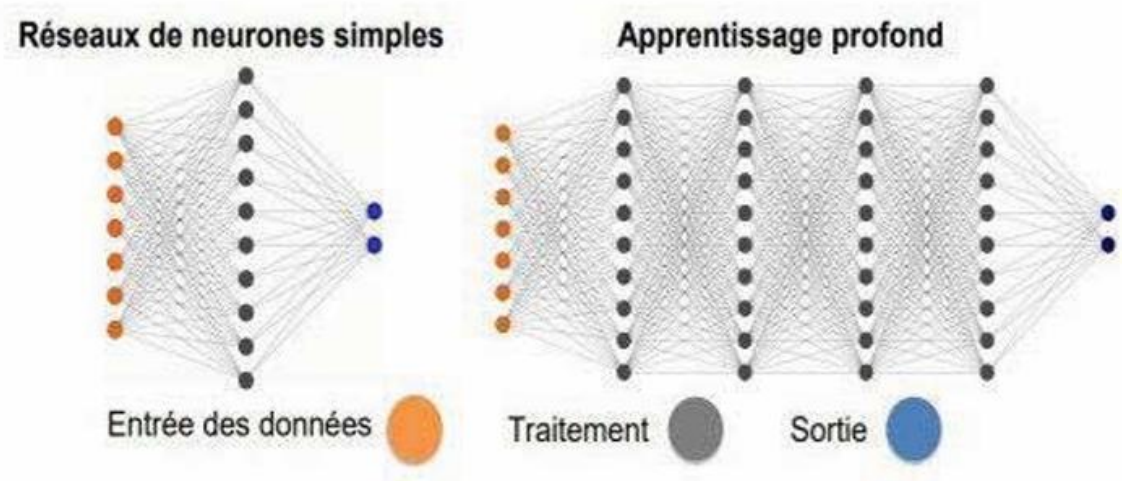


Figure III-4 : l'architecture d'un Réseau de neurones

Biological neurons	Neurones artificiels
Synapses	Connections pondérés
Axons	Sorties
Dendrites	Entrées
Summinator	Fonction d'activation

Tableau III-1 : Analogie entre neurones biologiques et artificiels

Modèle mathématique d'un neurone artificiel :

Un neurone artificiel est un modèle mathématique qui prend plusieurs entrées, les pondère et les somme pour produire une sortie. Figure III.5 est la représentation classique d'un RNA [40] :

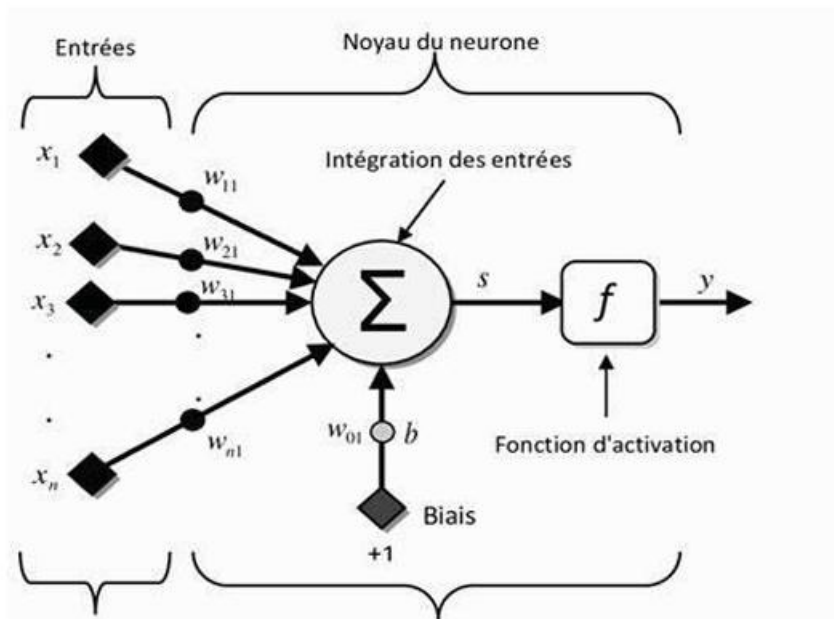


Figure III-5 Modèle d'un neurone artificiel (Marc Parizeau., 2004)

Ce modèle mathématique peut être représenté comme suit :

$$S = \sum_{j=1}^R x_j w_{1j} + b \quad (\text{III- 1})$$

Où :

s est la sortie du neurone

x_j Sont les entrées du neurone

$w_{1,j}$ sont les poids associés à chaque entrée

b est le biais (ou offset)

et :

$$y = f(s) = f\left(\sum_{j=1}^R x_j w_{1j} + b\right) \quad (\text{III- 2})$$

où f est une fonction d'activation qui transforme la sortie pondérée s en une sortie non-linéaire.

III.3 Apprentissage des réseaux neuronaux artificiels :

L'apprentissage des réseaux neuronaux artificiels consiste à ajuster les poids des connexions entre les neurones du réseau afin qu'il puisse apprendre à effectuer des tâches spécifiques. Le processus d'apprentissage implique généralement deux phases principales : la phase de propagation avant (forward propagation) et la phase de rétropropagation du gradient (backpropagation) [41].

Vocabulaire des Réseaux Neuronaux :

Il est vraiment important de bien connaître les concepts et le vocabulaire liés à l'apprentissage des réseaux neuronaux pour comprendre comment fonctionne l'apprentissage profond. En comprenant ces termes, on peut mieux saisir comment un réseau de neurones est construit et comment il apprend dans ce domaine.

III.4 -Les réseaux CNN (Convolutional Neural Networks)

Un réseau neuronal convolutif s'inspire de la nature, car la façon dont les neurones artificiels sont connectés rappelle l'organisation du cortex visuel chez l'animal.

L'un des usages principaux de ces technologies, c'est la reconnaissance d'images. Les réseaux convolutifs apprennent vraiment vite et ont tendance à faire moins d'erreurs. Moins souvent, on les utilise aussi pour analyser des vidéos.

Ce type de réseau est aussi utilisé pour le traitement naturel du langage. Les modèles CNN sont vraiment efficaces pour analyser la sémantique, modéliser des phrases, classer du contenu ou même pour la traduction.

Comparés aux méthodes plus traditionnelles comme les réseaux de neurones récurrents, les réseaux de neurones à convolution peuvent capter différents détails du contexte dans le langage, sans avoir besoin de suivre une séquence linéaire comme c'est le cas avec les modèles récurrents.

Les réseaux convolutifs sont aussi utilisés pour aider à découvrir de nouveaux médicaments. Ils peuvent prévoir comment différentes molécules interagissent avec des protéines dans le corps, ce qui facilite la recherche de traitements efficaces.

Les CNN ont notamment été utilisés pour les logiciels de jeu de Go ou de jeu d'échecs auxquels ils sont capables d'exceller. Une autre application est la détection d'anomalie sur une image en entrée (input). [40]

III.5 -Architecture d'un Convolutional Neural Network-CNN

Un CNN, ou Réseau de Neurones Convolutif, c'est un type de réseau conçu pour analyser des données structurées comme des images ou des signaux. Il fonctionne avec plusieurs couches spécialisées qui aident à repérer petit à petit les caractéristiques importantes à différentes échelles, ce qui facilite la reconnaissance et l'analyse des images ou autres données complexes.

Architecture d'un Convolutional Neural Network (CNN)

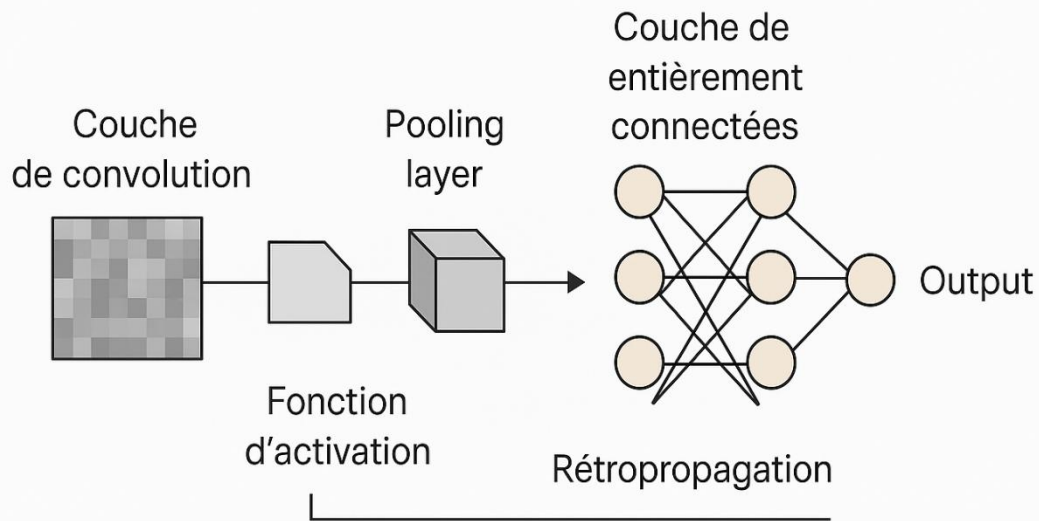


Figure III-6 : l'architecture d'un CNN

La figure illustre l'architecture d'un réseau neuronal convolutif (CNN) et le flux des données à travers ses différentes couches :

- ❖ **Couche d'entrée** : Reçoit les données brutes sous forme de matrices (images, signaux, etc.).
- ❖ **Couches de convolution** : Appliquent des filtres pour détecter des motifs et extraire des caractéristiques essentielles.
- ❖ **Couches de pooling** : Réduisent la taille des données tout en conservant les informations clés, améliorant ainsi l'efficacité du traitement.
- ❖ **Couches entièrement connectées** : Analysent les caractéristiques extraites et réalisent les prises de décision.
- ❖ **Couche de sortie** : Génère la prédiction finale du modèle (classification, détection, etc.).

Cette architecture est largement utilisée pour le traitement d'images, la reconnaissance vocale et la suppression du bruit dans des systèmes comme l'OFDM.

III.6 L'auto encodeur

Un auto encodeur est un type de réseau neuronal conçu pour apprendre à reproduire fidèlement son entrée tout en extrayant une représentation compacte des données. Sa structure comprend une couche d'entrée et une couche de sortie ayant le même nombre de neurones.

Son fonctionnement repose sur deux étapes clés :

- ❖ **Encodage** : L'autoencodeur réduit la dimensionnalité des données en les transformant en une représentation latente plus condensée.

- ❖ **Décodage** : Il reconstruit les données à partir de cette représentation en minimisant la perte d'information.

L'architecture d'un auto-encodeur :

L'auto encodeur est composée de deux parties principales : un encodeur f_θ et un décodeur g_ψ . L'encodeur calcule $z_i = f_\theta(x_i)$ pour chaque échantillon d'apprentissage en entrée (x_i) avec $i \in [1, n]$; et $i \in [1, p]$, le décodeur vise à reconstituer l'entité à partir de code z_i : $\hat{x}_i = g_\psi(z_i)$

Les paramètres de l'encodeur et du décodeur sont ajustés simultanément lors de la tâche de reconstruction, en cherchant à minimiser la fonction objective.

$$\mathcal{L}_{AE}(\theta, \psi) = \mathcal{L}(x_i, g_\psi(f_\theta(x_i))) = \sum_{i=1}^n \mathcal{L}(x_i, g_\psi(z_i)) = \sum_{i=1}^n \mathcal{L}(x_i, \hat{x}_i) \quad (\text{III- 3})$$

-L'encodeur

L'encodeur joue un rôle essentiel dans la compression des données d'entrée en une représentation plus concise. En effectuant cette compression, il identifie les caractéristiques les plus significatives à partir des données initiales, ce qui conduit à la création d'une représentation condensée appelée "bottleneck" ou espace latent.

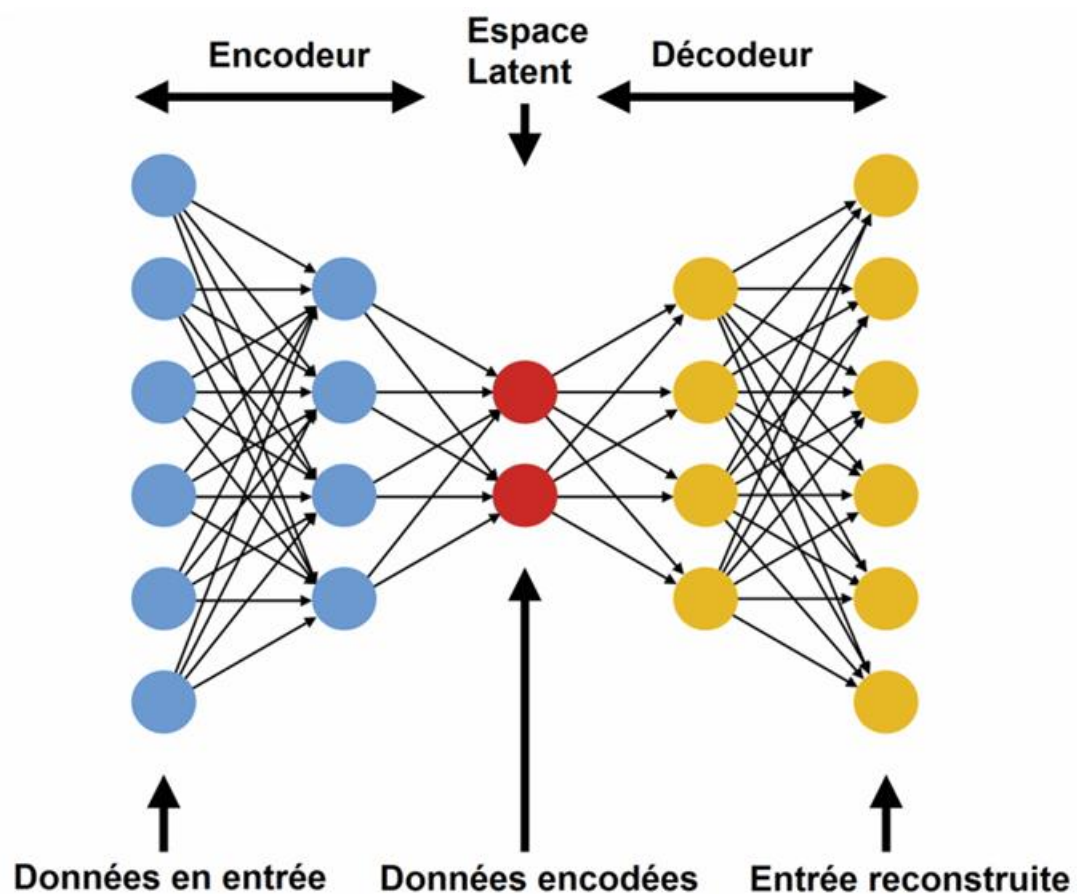


Figure III - 7 : L'architecture d'un auto-encodeur

-Le décodeur

Contrairement à l'encodeur, le décodeur sert à décompresser le message compressé pour retrouver les données originales. Son principal défi, c'est d'utiliser au mieux les informations contenues dans le vecteur réduit pour recréer le fichier ou l'image autant que possible, et que ce soit fidèle à l'original

L'espace latent

L'espace latent, c'est une sorte de résumé compact des données, placé entre l'encodeur et le décodeur. En gros, il sert à filtrer tout ce qui n'est pas essentiel, pour ne garder que l'essence des informations. Ça permet d'éliminer le bruit et de rendre le modèle plus efficace et plus précis.

Paramètres d'entraînement pour un auto-encodeur :

Lors de l'entraînement d'un autoencodeur, l'algorithme de descente de gradient est souvent utilisé pour optimiser les performances du modèle. Plusieurs paramètres influencent la qualité de l'apprentissage :

- ❖ **Taille du code (couche cachée)** : Détermine le nombre de neurones dans la représentation latente, influençant le niveau de compression des données.
- ❖ **Nombre de couches** : Plus un autoencodeur a de couches, plus il peut capturer des caractéristiques complexes des données.
- ❖ **Nombre de neurones par couche** : L'encodeur réduit progressivement la dimensionnalité, tandis que le décodeur la reconstitue, formant une structure symétrique.
- ❖ **Fonction objective** : L'erreur quadratique moyenne est souvent utilisée, mais pour des données probabilistes (entre 0 et 1), l'entropie croisée est plus adaptée.
- ❖ **Nombre d'époques** : Détermine combien de fois le modèle est entraîné sur l'ensemble des données. Un nombre trop faible entraîne un sous-apprentissage, tandis qu'un nombre trop élevé peut conduire au sur-apprentissage.

En ajustant ces paramètres correctement, un autoencodeur peut optimiser la reconstruction des données et améliorer ses performances pour des tâches comme la réduction de dimension ou la détection d'anomalies.

Principe de fonctionnement d'un auto-encodeur :

Le principe de fonctionnement de la méthode auto encodeur est le suivant :

Structure d'un autoencodeur

Un autoencodeur est un réseau de neurones artificiels utilisé en apprentissage non supervisé. Il est constitué de deux parties principales :

- ❖ **L'encodeur** : Il prend une donnée d'entrée et la transforme en une représentation latente compressée.
- ❖ **Le décodeur** : Il reconstruit cette représentation en essayant de retrouver la donnée d'origine.

Ces deux parties sont reliées par une couche centrale, appelée code latent, qui stocke une version condensée des informations essentielles

.Étapes du fonctionnement

- ❖ **Encodage.**

L'encodeur reçoit une donnée, comme une image, un texte ou un signal. Ensuite, il utilise plusieurs couches de neurones pour transformer cette donnée, en réduisant sa taille tout en conservant ses éléments importants.

❖ **Stockage du code latent :**

Après l'encodage, on transforme les données en une version plus compacte, qu'on appelle un code latent. Ce code contient uniquement ce qu'il faut pour reconstruire l'information, rien de superflu.

❖ **Décodage :**

Après l'encodage, on transforme les données en une version plus compacte, qu'on appelle un code latent. Ce code contient uniquement ce qu'il faut pour reconstruire l'information, rien de superflu.

❖ **Comparaison et ajustement :**

L'autoencodeur compare la donnée qu'il a reconstruit avec l'entrée originale et ajuste ses poids neuronaux pour réduire l'écart entre les deux. Pour ce faire, il utilise un processus d'optimisation, le plus souvent basé sur la technique appelée descente de gradient.

Optimisation du modèle

L'entraînement de l'autoencodeur vise à améliorer la précision de reconstruction en ajustant plusieurs paramètres

- ❖ **Taille du code latent :** Plus la représentation latente est petite, plus l'autoencodeur compresse l'information.
- ❖ **Nombre de couches :** Un autoencodeur peut être simple (une seule couche) ou profond (plusieurs couches) selon la complexité des données à traiter.
- ❖ **Nombre de neurones par couche :** L'encodeur réduit progressivement le nombre de neurones, tandis que le décodeur les augmente symétriquement.
- ❖ **Fonction de perte :** L'erreur de reconstruction est mesurée par l'erreur quadratique moyenne ou l'entropie croisée, en fonction du type de données.
- ❖ **Nombre d'époques :** Il définit combien de fois le modèle est exposé aux données d'entraînement pour affiner ses poids neuronaux.

Applications d'un autoencodeur

L'autoencodeur est un outil puissant utilisé dans plusieurs domaines :

- ❖ **Réduction de dimension :** Il permet de simplifier des jeux de données volumineux sans perdre leur essence.
- ❖ **Détection d'anomalies :** Il repère des données aberrantes en comparant l'entrée et la sortie.
- ❖ **Compression de données :** Il réduit la taille des fichiers en conservant uniquement les éléments essentiels.

- ❖ **Génération de données** : Des variantes comme les autoencodeurs variationnels (VAE) sont capables de créer de nouvelles données synthétiques

Exemple concret

Prenons l'exemple d'un autoencodeur utilisé pour nettoyer des images :

- On lui fournit une image bruitée (avec des pixels altérés).
- L'encodeur extrait une version propre de l'image en supprimant le bruit.
- Le décodeur reconstruit l'image sans les imperfections initiales.
- Le modèle ajuste ses paramètres pour améliorer la qualité du nettoyage sur de nouvelles images

L'autoencodeur est une architecture très importante dans le domaine de l'intelligence artificielle.

Elle est souvent utilisée pour analyser, compresser et reconstruire des données de façon efficace. On la retrouve notamment en vision par ordinateur, en reconnaissance de motifs, et dans le traitement du langage naturel.

III.7 DENOISING AUTO-ENCODER (DAE) DANS UNE CHAÎNE OFDM

La transmission de signaux en OFDM peut rencontrer différents types de bruit, comme le bruit blanc additif (AWGN) ou le bruit impulsif, ce qui peut sérieusement affecter la qualité de la communication. Pour faire face à ces problèmes, l'utilisation d'auto-encodeurs débruiteurs (DAE) semble très prometteuse. Dans ce mémoire, on va explorer comment concevoir et évaluer un DAE capable d'atténuer le bruit impulsif dans une chaîne OFDM.

. Génération de signaux bruités – Partie émission :

Pour évaluer l'efficacité du DAE, nous commençons par générer des données OFDM simulées. Ces données suivent les étapes de transmission classiques :

Modulation numérique : QPSK ou 16-QAM pour encoder les symboles.

Transformation par IFFT : Permet de passer du domaine fréquentiel au domaine temporel.

Ajout du préfixe cyclique (CP) : Assure une meilleure gestion de l'interférence inter-symbole.

Modélisation du bruit

- ❖ **Bruit Gaussien Additif (AWGN)** : Modélise les perturbations thermiques classiques.

Bruit impulsif : Utilisation de modèles comme Bernoulli-Gaussian ou Middleton Class A pour représenter des perturbations non gaussiennes souvent rencontrées en transmission radio.



Résultats de
simulation

IV.1 Introduction

Dans ce chapitre, on présente en premier lieu les résultats d'une simulation d'une chaîne de communication OFDM dans laquelle on a injecté un bruit blanc gaussien et en second lieu les résultats de l'application des techniques de deep Learning, notamment l'autoencodage pour la réduction d'un bruit impulsif injecté artificiellement après la phase d'encodage. On commence par expliquer brièvement pourquoi on a choisi l'environnement de développement Google Colab pour réaliser cette simulation en langage Python, en soulignant qu'il est parfaitement adapté pour modéliser des systèmes de communications complexes. Les résultats obtenus de cette simulation seront présentés sous forme de graphiques et figures, permettant ainsi d'évaluer les performances du système qu'on a simulé, comme la vitesse de transmission et la résistance aux interférences et au bruit.

IV.2 Présentation de Google colab

Google Colaboratory, ou "Colab", est une plateforme gratuite basée sur Jupyter Notebooks, permettant aux utilisateurs de coder en Python directement dans leur navigateur sans configuration préalable. Destinée à l'enseignement et à la recherche en IA et en data science, elle offre un accès gratuit à des GPU pour accélérer le calcul et l'exécution des projets de machine learning. Les notebooks peuvent être partagés facilement, ce qui favorise la collaboration et la reproduction des travaux. En somme, Colab est un outil puissant et accessible pour l'analyse de données, le développement et l'apprentissage machine.



Figure IV-1 : the collaboriez

IV.3 Pourquoi choisir Google Colab ?

Google Colaboratory, ou "Colab", présente de nombreux avantages :

- ❖ **Gratuité** : Accès gratuit aux ressources, y compris GPU et TPU, permettant d'accélérer les calculs et l'entraînement des modèles de machine learning.
- ❖ **Accessibilité** : Fonctionne directement dans le navigateur sans nécessiter d'installation ou de configuration préalable.
- ❖ **Partage facile** : Notebooks facilement partageables avec des collègues ou étudiants, favorisant la collaboration et la reproduction des travaux.

- ❖ . Intégration avec Google Drive : Permet de sauvegarder et de synchroniser les notebooks avec Google Drive pour une accessibilité simplifiée.
- ❖ Support des bibliothèques populaires : Compatible avec de nombreuses bibliothèques Python couramment utilisées en data science et machine Learning (TensorFlow, Keras, PyTorch, etc.).

IV.4 Simulation et Résultats

Les signaux sont toujours affectés par du bruit, qu'il provienne des appareils de mesure ou du support de transmission. Pour le réduire, on utilise des techniques mathématiques comme la soustraction de la moyenne. Cependant, ces méthodes supposent que l'on connaît la nature du bruit, ce qui n'est pas toujours le cas, notamment pour le bruit impulsif. Dans ces situations, une approche plus efficace consiste à apprendre à identifier le bruit à partir de données d'exemple, ce qui sera exploré dans notre simulation.

Simulation de la chaîne OFDM

Dans cette section, nous avons effectué une simulation de la chaîne OFDM dans un canal de Rayleigh avec un bruit blanc gaussien. La simulation a été réalisée pour la bande de fréquence [27,5-28,35] GHz, qui a été attribuée et licenciée par les États-Unis pour la cinquième génération. Les paramètres utilisés sont présentés dans le tableau ci-dessous :

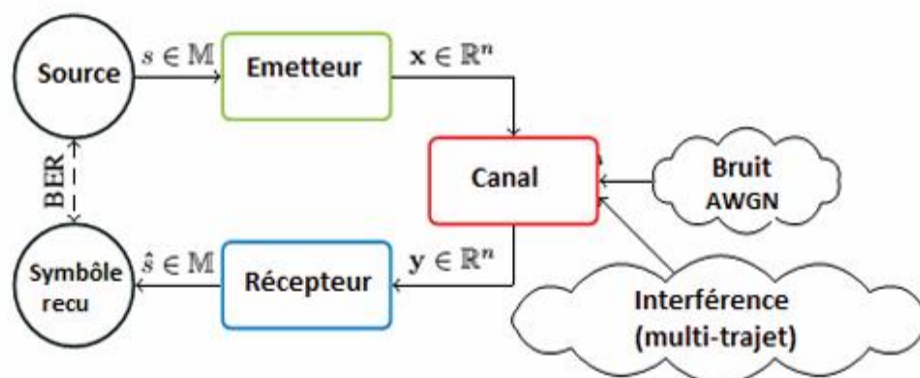


Figure IV-2 Schéma bloc de la chaîne de transmission

Paramètres	Valeurs
Technique de modulation	BPSK, QPSK, M-QAM
Nombres de sous-porteuses	512
Longueur de préfix cyclique	128
Nombres d'FFT	512
Bande de fréquence	[27,5-28,38] GHz
Modèle de Canal	Canal de Rayleigh

Tableau IV-1 Paramètres de simulation

Dans ce qui suit nous avons s'intéressés par les performances en termes de taux d'erreur en fonction de rapport Signal sur Bruit.

Le taux d'erreur binaire (TEB ou BER)

Le terme BER (Bit Error Rate) ou TEB (Taux d'Erreur Binaire) est utilisé pour quantifier les erreurs qui surviennent dans un système de transmission. Il s'agit d'un paramètre essentiel pour évaluer les systèmes qui assurent le transfert de données numériques d'un emplacement à un autre [45]. La formule suivante définit le taux d'erreur binaire :

$$BER = TER = \frac{\text{Nombre de bits erronés}}{\text{Nombre de bits totals envoyés}} \quad (IV- 1)$$

Le rapport signal sur bruit (SNR)

SNR (Signal to Noise Ratio) est une mesure du rapport signal sur bruit, qui correspond à la division de la puissance du signal utile (P_s) par la puissance du bruit (N). Cette mesure permet d'évaluer la contribution du bruit et son impact sur la dégradation du signal.

$$SNR = \frac{p_s}{n_0} \quad (IV- 2)$$

Ce rapport est généralement exprimé en dB :

$$SNR = 10 \log \left(\frac{p_s}{n_0} \right) \quad (IV- 3)$$

Le SNR (Signal to Noise Ratio) est une mesure essentielle permettant d'évaluer la dégradation soudaine causée par le bruit. En effet, lorsque le rapport signal sur bruit est faible, le signal est davantage affecté par le bruit, rendant ainsi plus difficile l'élimination de son influence. Pour garantir l'intégrité du signal reçu en tant que "copie fidèle" du signal transmis, il est essentiel de maintenir un rapport signal sur bruit élevé [46].

La figure IV.3 représente le BER en fonction de SNR pour un bruit blanc AWGN en utilisant la modulation 16-QAM on remarque que la relation entre le BER et SNR est nettement inversement proportionnelle, plus le SNR augment plus le BER diminue.

IV.5 PARTIE 01 : simulation avec bruit Gaussien

Analyse de la Densité Spectrale de Puissance (DSP) et du Taux d'Erreur Binaire (BER) pour les Modulations BPSK, QPSK et 16-QAM en présence de Bruit Gaussien

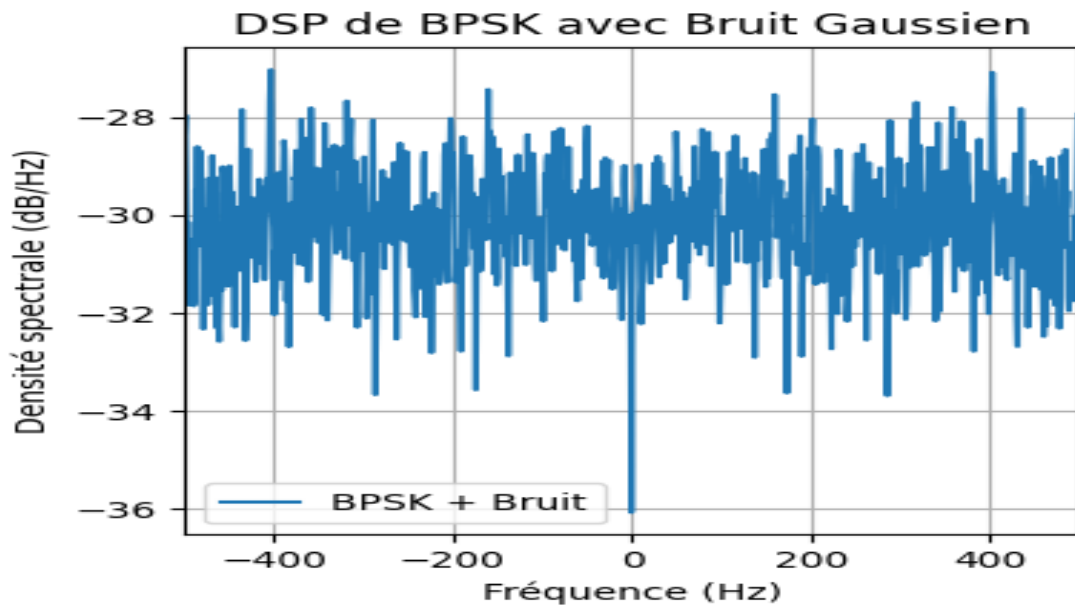


Figure IV-3 : DSP de BPSK avec bruit Gaussien

Ce graphique illustre l'impact du bruit sur la représentation dans le domaine fréquentiel d'un signal de communication numérique. Le bruit introduit un caractère aléatoire sur toutes les fréquences, ce qui rend plus difficile la distinction claire des caractéristiques spectrales du signal BPSK original. Pour mieux analyser le signal BPSK lui-même, on pourrait envisager des techniques pour réduire le bruit, telles que le filtrage ou le moyennage de plusieurs estimations spectrales

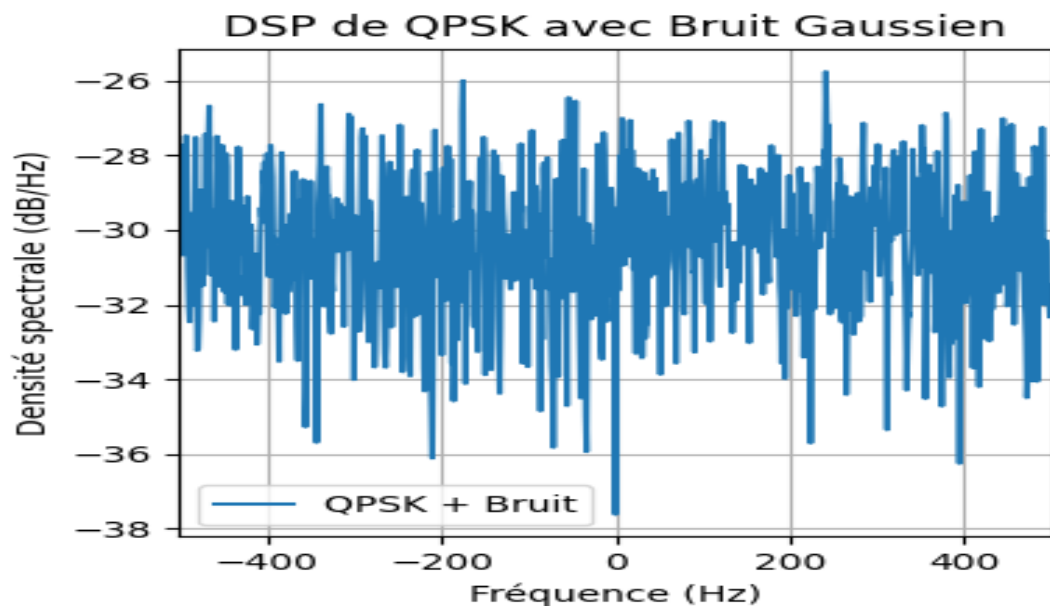


Figure IV-4 : DSP de QPSK avec bruit Gaussien

Cette figure met en évidence l'effet du bruit sur le spectre d'un signal QPSK. Le bruit gaussien, avec sa distribution de puissance sur une large bande de fréquences, rend difficile l'identification des caractéristiques spectrales propres au signal QPSK. Pour une analyse plus approfondie du signal QPSK seul, des techniques de réduction du bruit seraient nécessaires.

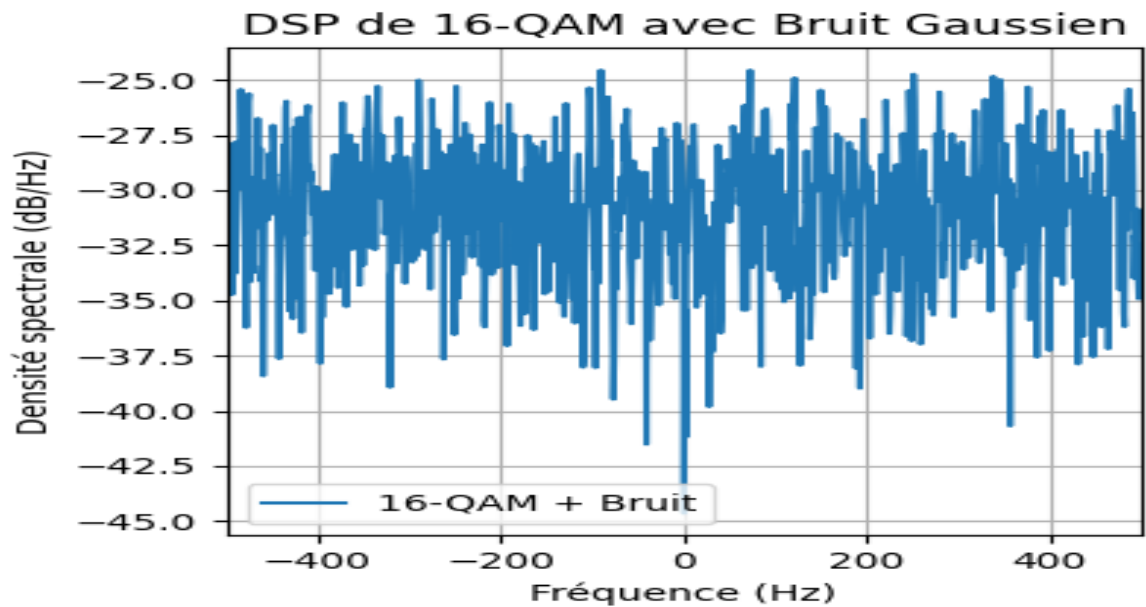


Figure IV-5 : DSP de 16-QAM avec bruit Gaussien

Cette figure montre comment le bruit gaussien peut obscurcir les caractéristiques spectrales d'un signal de modulation d'ordre supérieur comme le 16-QAM. La présence du bruit rend difficile l'identification précise de la forme spectrale intrinsèque du signal 16-QAM. Pour une analyse plus détaillée du signal 16-QAM seul, des techniques de réduction du bruit seraient nécessaires afin de révéler sa signature spectrale sous-jacente.

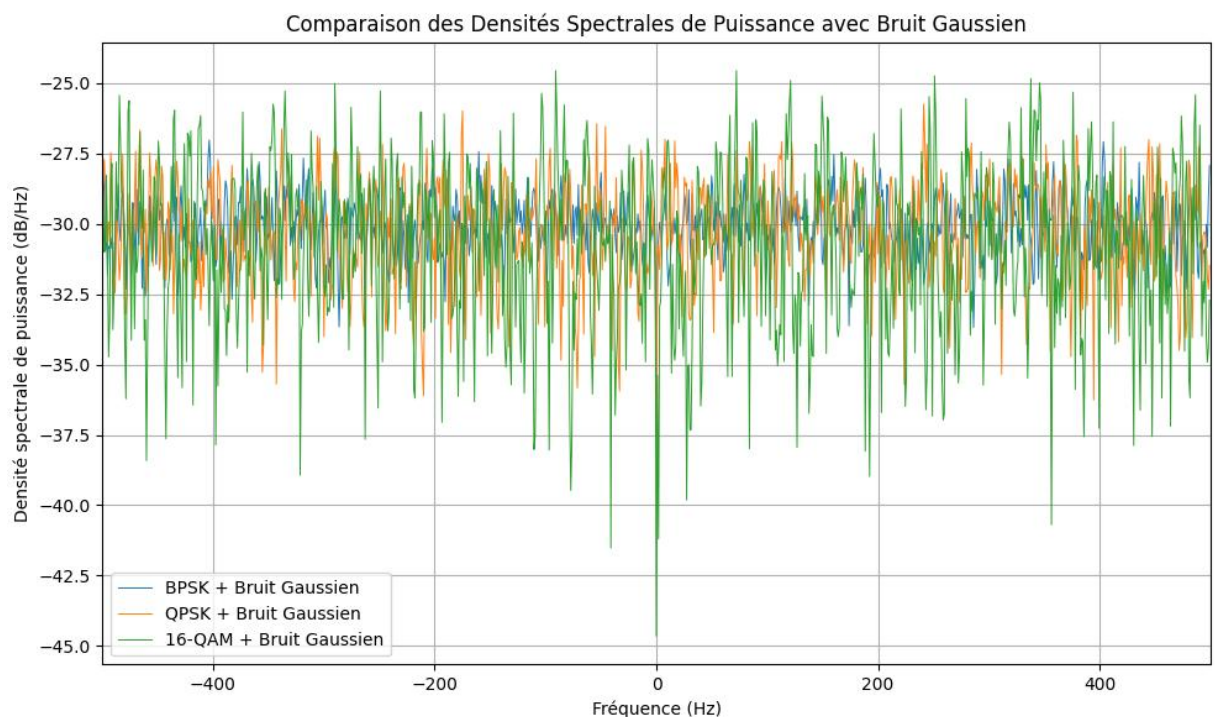


Figure IV-6 : la comparaison des densités spectrales de puissance avec bruit Gaussien

Cette figure met en évidence l'impact du bruit gaussien sur le spectre de différents signaux modulés et souligne la difficulté de distinguer leurs caractéristiques spectrales intrinsèques lorsque le bruit est significatif. Bien que les types de modulation aient des efficacités spectrales différentes en théorie, leur représentation dans le domaine fréquentiel est ici fortement influencée par le bruit, rendant leur différenciation visuelle complexe. Pour une analyse plus approfondie des différences spectrales entre ces modulations, il serait nécessaire d'observer les signaux avec un rapport signal sur bruit (SNR) plus élevé ou d'appliquer des techniques de traitement du signal pour atténuer l'effet du bruit.

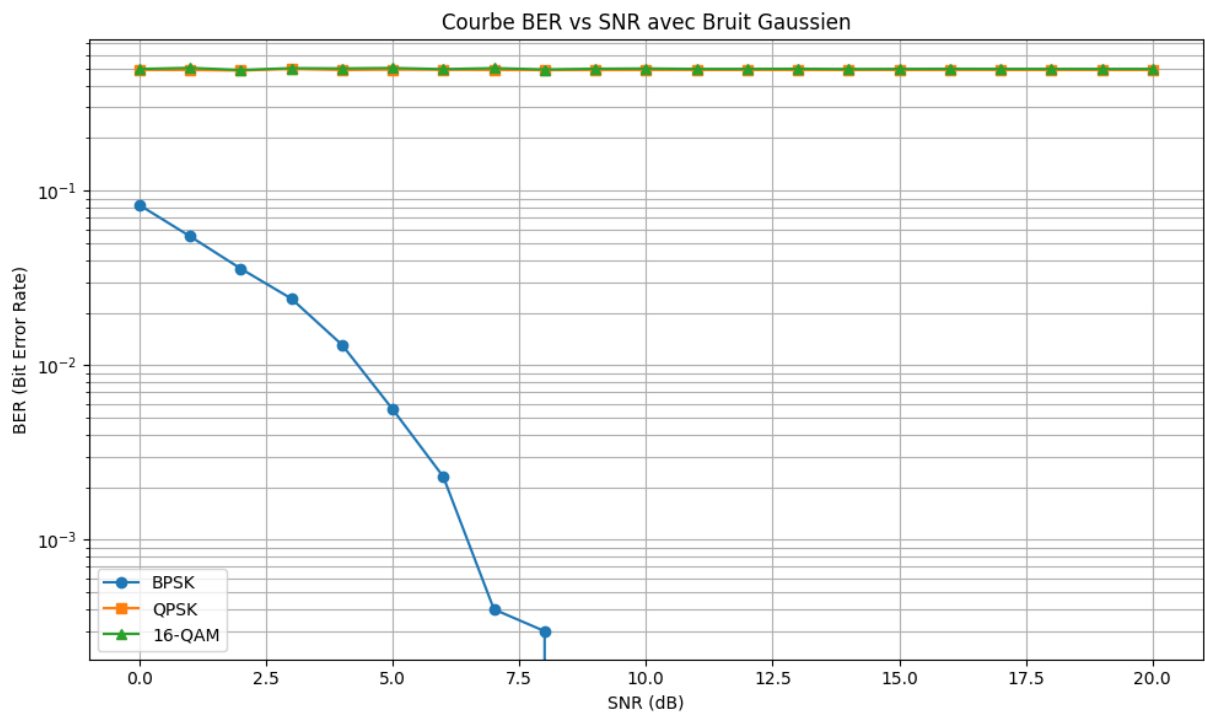


Figure IV-7 : la courbe de BER vs SNR avec bruit gaussien.

Cette figure compare l'efficacité de différentes techniques de modulation en présence de bruit gaussien en termes de compromis entre le rapport signal sur bruit et le taux d'erreur binaire. Elle illustre clairement que pour une même puissance de signal relative au bruit, les modulations d'ordre supérieur comme le 16-QAM ont une performance en termes de BER supérieure à celles d'ordre inférieur comme le BPSK et le QPSK. Pour obtenir une performance similaire, les modulations d'ordre supérieur nécessitent un SNR plus élevé, ce qui reflète leur plus grande efficacité spectrale (plus de bits transmis par symbole) au prix d'une plus grande sensibilité au bruit.

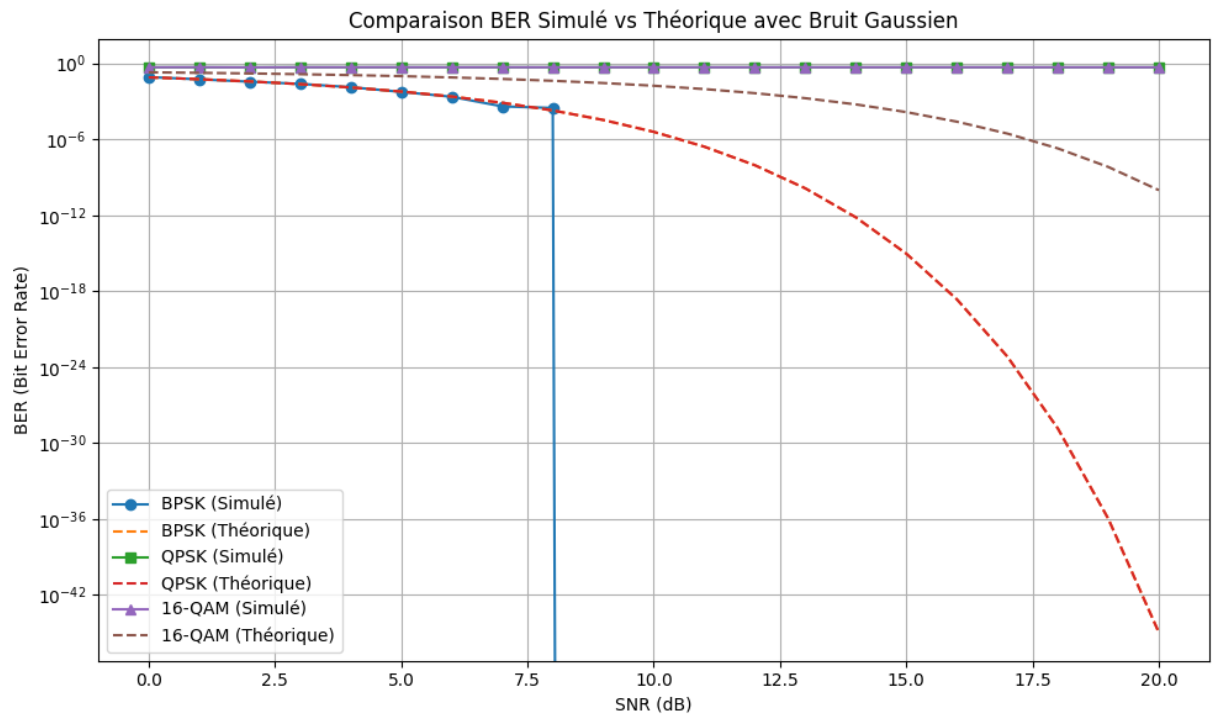


Figure IV-8 : la comparaison BER simulé vs théorique avec bruit Gaussien

Cette figure est très instructive car elle permet de valider la fidélité des simulations par rapport aux prédictions théoriques pour la performance des systèmes de communication numériques en présence de bruit gaussien. La bonne concordance pour le BPSK et le QPSK renforce la confiance dans les modèles théoriques. L'écart observé pour le 16-QAM souligne l'importance de considérer les imperfections et les approximations inhérentes aux simulations, et potentiellement la nécessité d'affiner les modèles théoriques pour les modulations d'ordre supérieur dans des scénarios réels

Conclusion

A travers cette étude, nous avons analysé la performance des modulations BPSK, QPSK et 16-QAM en simulant la présence du bruit gaussien. Il a été prouvé que les modulations d'ordre supérieur (comme 16-QAM) offrent une meilleure efficacité spectrale mais sont plus sensibles au bruit, nécessitant un SNR plus élevé pour un même TEB. Inversement, BPSK et QPSK sont plus robustes au bruit mais moins efficaces spectralement. L'analyse spectrale a confirmé l'impact du bruit qui masque les caractéristiques du signal. La validité des simulations a été confirmée pour BPSK et QPSK, soulignant l'importance des considérations pratiques pour les modulations plus complexes. En conclusion, maintenir un bon SNR est crucial, et un compromis entre efficacité spectrale et robustesse au bruit est inévitable.

IV.6 Partie 02 : simulation avec bruit impulsif

. COMMENT IMPLEMENTER UN AUTO ENCODEUR POUR LE DEBRUITAGE D'UN SIGNAL (DENOISING AUTO-ENCODEUR DAE)

On pourrait implémenter ça en Python, en utilisant une librairie populaire comme TensorFlow ou PyTorch pour la partie réseau de neurones.

On pourrait implémenter ça en Python, en utilisant une librairie populaire comme TensorFlow ou PyTorch pour la partie réseau de neurones.

Voici une façon de procéder en utilisant TensorFlow et Keras :

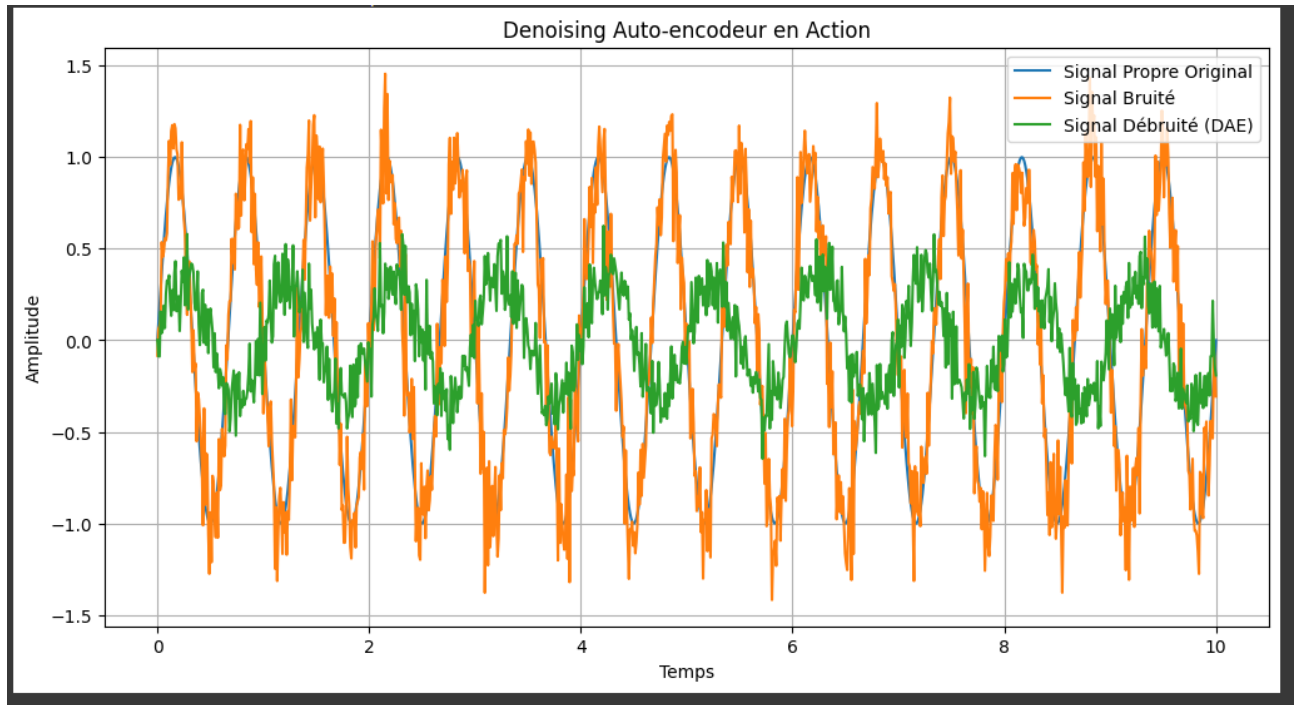


Figure IV-9 : Débruitage Auto-encodeur en action

La figure illustre l'effet d'un auto-encodeur de débruitage (DAE) sur un signal bruité. On observe trois courbes :

- ❖ **Signal propre original (bleu) :** Il représente le signal sans perturbation, servant de référence.
- ❖ **Signal bruité (orange) :** Ce signal a été altéré par du bruit, ce qui complique son exploitation.
- ❖ **Signal débruité par le DAE (vert) :** Après traitement par l'auto-encodeur, le signal retrouve une forme proche de l'original.

L'axe des abscisses représente le temps, tandis que l'axe des ordonnées indique l'amplitude du signal. L'amélioration apportée par le DAE est visible, car le signal débruité est plus proche du signal original que du signal bruité, montrant l'efficacité du modèle pour atténuer les perturbations.

Points importants à considérer pour votre application spécifique (par exemple, le débruitage de signaux OFDM) :

- ❖ **Le type de bruit :** Assurez-vous que la fonction ajouter_bruit simule de manière réaliste le type de bruit présent dans vos signaux OFDM.
- ❖ **L'architecture de l'auto-encodeur :** L'architecture (nombre de couches, taille des couches, fonctions d'activation) peut avoir un impact significatif sur les performances. Vous pourriez expérimenter avec des architectures plus profondes ou différentes. Pour des signaux temporels comme l'OFDM, des

couches récurrentes (RNN) ou des convolutions 1D pourraient également être envisagées dans l'encodeur et le décodeur pour capturer les dépendances temporelles.

- ❖ **La taille du code (taille_code)** : Ce paramètre contrôle la compression. Une taille de code plus petite peut forcer le réseau à apprendre des caractéristiques plus importantes, mais une taille trop petite peut entraîner une perte d'informations importantes.
- ❖ **La fonction de perte** : mse est un bon point de départ, mais d'autres fonctions de perte pourraient être plus appropriées en fonction des caractéristiques de vos données et du type de bruit.
- ❖ **La normalisation des données** : Il est souvent bénéfique de normaliser vos données (par exemple, en les mettant à l'échelle entre 0 et 1 ou en les standardisant avec une moyenne de 0 et un écart-type de 1) avant de les entraîner. Cela peut aider à la convergence de l'entraînement.
- ❖ **L'évaluation** : Après l'entraînement, il est crucial d'évaluer les performances de votre DAE sur un ensemble de données de test indépendant pour voir s'il généralise bien et s'il atteint le niveau de débruitage souhaité. Vous pouvez utiliser des métriques comme le rapport signal sur bruit (SNR) pour quantifier l'amélioration.

En adaptant ces éléments à vos données OFDM spécifiques, vous pourrez construire un DAE personnalisé efficace pour le débruitage

Implémentation de l'auto-encodeur dans la chaîne OFDM

Dans les environnements réels, la suppression du bruit est complexe en raison de la diversité des interférences. Les approches classiques ne suffisent pas à traiter le bruit impulsionnel. Pour y remédier, nous introduisons un modèle d'apprentissage profond capable d'atténuer ces perturbations efficacement.

Un DAE (Denoising Auto-encoder) est un modèle de débruitage qui apprend à supprimer les perturbations spécifiques à un type de données. Lorsqu'il est entraîné sur un ensemble particulier, il optimise la suppression du bruit pour des données similaires. Par exemple, un DAE entraîné sur des signaux OFDM sera efficace pour ces signaux mais inadapté au traitement d'images.

Pour l'implémentation de l'auto-encodeur nous avons suivi le modèle trouvé dans [43] en introduisons le bloc AE avant la modulation OFDM figure IV.10.

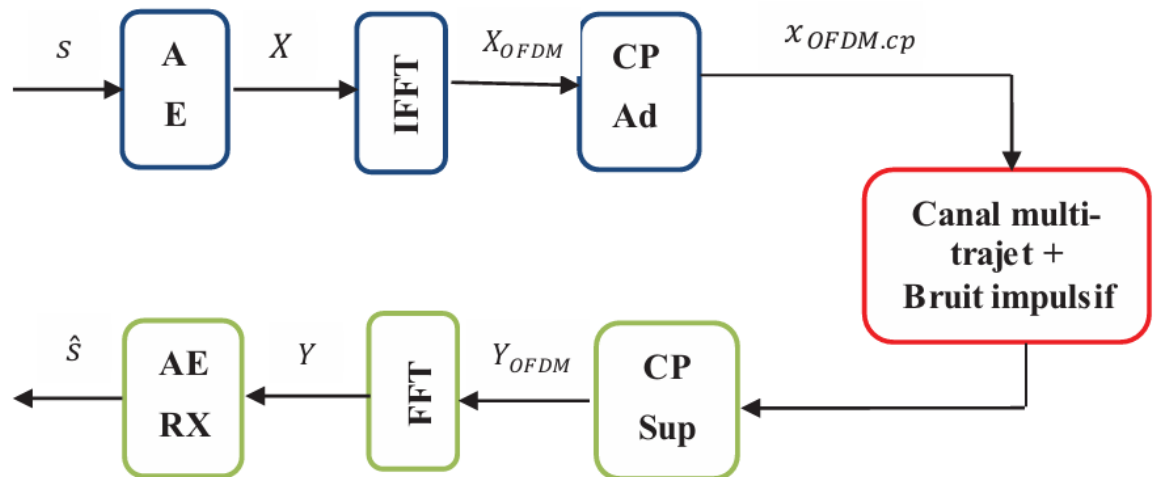


Figure IV-10 : Schéma fonctionnel

L'architecture spécifique en tête, intégrant l'auto-encodeur (AE) avant la modulation OFDM à l'émission et un bloc de décompression AE (AE RX) après la démodulation FFT au récepteur.

C'est une approche intéressante pour potentiellement améliorer la robustesse de la transmission OFDM face au bruit, y compris le bruit impulsionnel mentionné.

Voici une explication de l'implémentation; implémentation en Python en suivant ce schéma, en utilisant

toujours TensorFlow/Keras comme exemple :

À l'Émetteur :

1. Bloc AE (Encodeur) :

- ❖ L'entrée de ce bloc est votre flux de données source s . Il pourrait s'agir d'une séquence de bits ou de symboles.
- ❖ Ce bloc AE est en réalité la partie encodeur de votre auto-encodeur. Son rôle est d'apprendre une représentation compressée et potentiellement plus robuste de vos données d'entrée s . La sortie de cet encodeur est X .
- ❖ **Implémentation Python** : Vous définiriez un modèle Keras qui prend en entrée la forme de vos données s et produit une sortie X de dimension inférieure (ou égale, mais avec une représentation différente).
- ❖ La sortie de l'encodeur AE, X , est ensuite passée à un bloc IFFT. Dans le contexte de l'OFDM, X représente les symboles modulés en fréquence. L'IFFT convertit ces symboles en un signal temporel X_{OFDM} .
- ❖ **Implémentation Python** : Vous utiliserez la fonction `np.fft.ifft()` de la librairie NumPy pour effectuer l'IFFT. Assurez-vous de gérer correctement la structure des données X (séquence de sous-porteuses).

3. Bloc CP Ad (Cyclic Prefix Addition) à corriger

- ❖ Un préfixe cyclique (CP) est ajouté au signal temporel X_{OFDM} pour créer $X_{\text{OFDM_cp}}$. Le CP est une copie de la fin du symbole OFDM ajoutée au début, ce qui aide à lutter contre l'interférence inter-symboles (ISI) causée par les trajets multiples du canal.
- ❖ **Implémentation Python** : Vous implémenterez une fonction qui prend le signal temporel et ajoute une copie de sa fin au début.
- ❖ Canal Multi-trajet + Bruit Impulsif :

Le signal $X_{\text{OFDM_cp}}$ traverse le canal de communication, qui est caractérisé par des trajets multiples et du bruit impulsionnel (et potentiellement d'autres types de bruit). C'est ici que les distorsions se produisent.

Au Récepteur :

1. Bloc CP Sup (Cyclic Prefix Suppression) :

- ❖ Au récepteur, la première étape consiste à supprimer le préfixe cyclique du signal reçu. On obtient ainsi Y_{OFDM} .
- ❖ **Implémentation Python** : Vous implémenterez une fonction qui supprime les premiers échantillons correspondant à la longueur du CP.
- ❖ Le signal temporel sans CP, Y_{OFDM} , est ensuite transformé dans le domaine fréquentiel à l'aide de la FFT, produisant Y .
- ❖ **Implémentation Python** : Vous utiliserez la fonction `np.fft.fft()` de NumPy.

3. Bloc AE RX (Décodeur AE) :

- ❖ La sortie de la FFT, Y (les symboles reçus dans le domaine fréquentiel, potentiellement corrompus par le canal), est l'entrée du bloc AE RX.
- ❖ Ce bloc AE RX est la partie décodeur de votre auto-encodeur. Il est entraîné à prendre la représentation potentiellement bruitée Y (qui correspond à une version bruitée de la représentation X encodée à l'émetteur) et à la reconstruire en une estimation du signal source original, $\hat{s}$ (s chapeau).
- ❖ **Implémentation Python** : Vous définirez un modèle Keras qui prend en entrée la forme de Y (qui aura la même taille que la sortie de l'encodeur X dans le domaine fréquentiel) et produit une sortie $\hat{s}$ de la même taille que l'entrée de l'encodeur s .

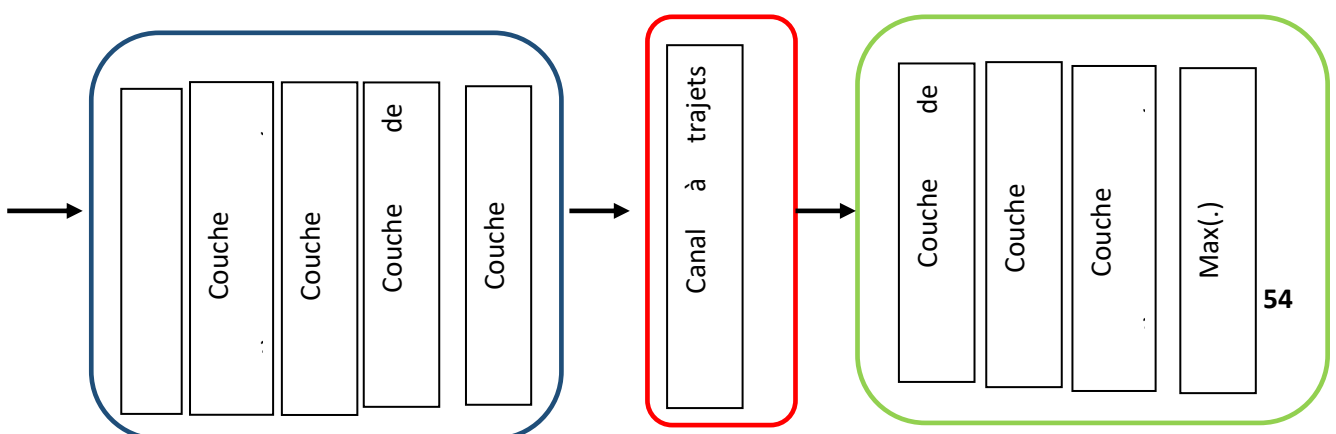


Figure IV-11 : Schéma fonctionnel d'un système de communication OFDM basé sur un Auto-Encodeur (AE)

Cette figure présente l'architecture d'un système de communication qui combine la modulation OFDM (Orthogonal Frequency Division Multiplexing) avec un auto-encodeur, une forme de réseau neuronal. L'objectif est de transmettre des informations de manière robuste à travers un canal de communication potentiellement bruyant et sujet à des trajets multiples. Le schéma est divisé en trois parties principales : l'émetteur, le canal de communication et le récepteur.

. L'Émetteur (Bloc Bleu) : Préparation du Signal pour la Transmission

L'émetteur est responsable de prendre le signal source et de le préparer pour une transmission efficace et fiable. Les étapes clés ici sont :

- ❖ **S (Signal d'Entrée)** : C'est le signal original que nous souhaitons transmettre. Il peut s'agir de données numériques représentées sous forme de vecteur, comme le suggère la colonne de valeurs.
- ❖ **Couche entièrement connectée avec ReLU** : Le signal d'entrée est d'abord traité par une couche de neurones entièrement connectée. Chaque neurone de cette couche est connecté à tous les éléments du signal d'entrée. L'application de la fonction d'activation ReLU (Rectified Linear Unit), définie par $\text{ReLU}(x) = \max(0, x)$, introduit une non-linéarité dans le modèle. Cette couche permet d'extraire des caractéristiques importantes du signal ou de créer une représentation plus appropriée.
- ❖ **Couche entièrement connectée avec activation linéaire** : Suite à la couche ReLU, une autre couche entièrement connectée est utilisée, mais cette fois avec une fonction d'activation linéaire. Cela signifie que la sortie de cette couche est une simple transformation linéaire de son entrée. Cette étape peut affiner davantage la représentation du signal avant la modulation.
- ❖ **Couche de normalisation** : Cette couche applique une technique de normalisation (comme la normalisation par lots ou la normalisation de couche) aux données. La normalisation aide à stabiliser l'entraînement du réseau neuronal et à améliorer sa capacité à généraliser en assurant une distribution cohérente des données.
- ❖ **Couche modulatrice OFDM** : C'est l'élément central pour la transmission. Cette couche prend la représentation du signal produite par le réseau neuronal et l'encode en utilisant la technique de modulation OFDM. L'OFDM divise la bande de fréquence disponible en de multiples sous-porteuses orthogonales, permettant de transmettre une partie du signal en parallèle sur chacune d'elles. Cela rend le système plus résistant aux interférences sélectives en fréquence et aux évanouissements dus aux trajets multiples.

- ❖ **X (Signal Transmis)** : Le résultat de la modulation OFDM est le signal transmis, noté 'X', qui est prêt à être envoyé à travers le canal de communication.

Le Canal (Bloc Rouge) : Introduction des Perturbations

- Cette figure présente l'architecture d'un système de communication qui combine la modulation OFDM (Orthogonal Frequency Division Multiplexing) avec un auto-encodeur, une forme de réseau neuronal. L'objectif est de transmettre des informations de manière robuste à travers un canal de communication potentiellement bruyant et sujet à des trajets multiples. Le schéma est divisé en trois parties principales : l'émetteur, le canal de communication et le récepteur.
- **Canal à trajets multiples Bruit implicite** : Cette description indique que le signal transmis 'X' atteint le récepteur par plusieurs chemins (trajets multiples) en raison de réflexions, de réfractions et de diffractions. Chaque trajet a un délai, une atténuation et un déphasage différents, ce qui déforme le signal. De plus, le canal introduit du "bruit implicite", qui est un signal aléatoire indésirable qui corrompt le signal transmis.
- **Y (Signal Reçu)** : Le signal qui arrive au récepteur, noté 'Y', est donc une version dégradée et bruitée du signal transmis 'X'.

Le Récepteur (Bloc Vert) : Récupération du Signal Original

Le récepteur a pour tâche de traiter le signal reçu 'Y' pour tenter de récupérer le signal d'origine 'S'.

Les étapes principales sont :

- **Couche de démodulateur OFDM** : La première étape au récepteur est la démodulation OFDM. Cette couche effectue l'opération inverse de la modulation OFDM, séparant le signal reçu 'Y' en les signaux portés par les différentes sous-porteuses. Elle annule l'effet du multiplexage en fréquence réalisé à l'émetteur.
- **Couche entièrement connectée avec ReLU** : Similaire à l'émetteur, le signal démodulé est ensuite traité par une couche entièrement connectée avec une activation ReLU. Cette couche commence probablement le processus d'extraction d'informations significatives du signal reçu et de compensation des distorsions introduites par le canal.
- **Couche entièrement connectée avec activation Softmax** : Une autre couche entièrement connectée suit, mais cette fois avec une fonction d'activation Softmax. La fonction Softmax est couramment utilisée dans la couche de sortie des réseaux de neurones pour la classification. Elle transforme un vecteur de nombres réels en une distribution de probabilités sur plusieurs classes possibles. Ici, cela suggère que l'auto-encodeur est conçu pour une tâche de classification, où l'objectif est de prédire l'entrée originale parmi un ensemble de catégories discrètes.

- **Max(-)** : Cela représente probablement une opération de maximum ou d'argmax. Si la sortie de la couche Softmax est une distribution de probabilités, cette opération sélectionne la classe ayant la probabilité la plus élevée comme étant la sortie reconstruite ou prédite.

En résumé, cette figure illustre un système de communication de bout en bout où un auto-encodeur basé sur un réseau neuronal est intégré à la modulation et à la démodulation OFDM. La partie "encodeur" de l'auto-encodeur (dans l'émetteur) apprend à créer une représentation efficace du signal d'entrée pour la transmission sur un canal bruyant et à trajets multiples. La partie "décodeur" de l'auto-encodeur (dans le récepteur) apprend à reconstruire le signal original à partir du signal reçu déformé.

IV.7 Entraînement de l'Auto-Encodeur :

L'étape cruciale est l'entraînement de l'auto-encodeur (l'encodeur AE et le décodeur AE RX). Vous devrez :

- ❖ **Générer des paires de données (entrée, sortie cible)** : L'entrée sera votre signal source s (ou une version encodée et passée par un canal simulé *sans* le bloc AE RX), et la sortie cible sera le même signal source s . L'objectif est que l'auto-encodeur apprenne à reconstruire l'entrée après avoir traversé l'encodeur, le canal (simulé avec le type de bruit attendu, y compris le bruit impulsionnel), et le décodeur.
- ❖ **Définir la fonction de perte et l'optimiseur** : Par exemple, l'erreur quadratique moyenne (MSE) pour des données continues ou l'entropie croisée pour des données catégorielles. Un optimiseur comme Adam est couramment utilisé.
- ❖ **Entraîner le modèle** : Utiliser la méthode fit de Keras pour entraîner l'auto-encodeur sur vos données.

IV.8 Utilisation pour la Transmission :

Une fois l'auto-encodeur entraîné :

- ❖ À l'émetteur, vous utilisez uniquement la partie encodeur (encodeur_ae) pour traiter vos données s avant la modulation OFDM.
- ❖ Au récepteur, après la démodulation OFDM (FFT), vous utilisez uniquement la partie décodeur (decodeur_ae) pour obtenir une estimation débruitée $\hat{s}$ des données originales.

Points Clés et Adaptations :

- **Forme des données** : Soyez très attentif à la forme de vos données à chaque étape (nombre d'échantillons, nombre de sous-porteuses pour l'OFDM, etc.) pour que les opérations matricielles (FFT, IFFT, les couches denses de l'AE) fonctionnent correctement.

- Représentation de Y : La nature exacte de Y (réel, complexe) déterminera la structure de l'entrée de votre décodeur AE. Si le canal introduit des distorsions de phase, votre AE devra peut-être apprendre à les compenser également.
- Entraînement end-to-end vs. séparé : Vous pourriez choisir d'entraîner l'AE séparément sur des données bruitées simulées, ou potentiellement d'envisager un entraînement "end-to-end" où l'AE est optimisé en tenant compte de l'ensemble du système de communication (ce qui est plus complexe).
- Complexité : L'ajout d'un AE introduit une complexité de calcul supplémentaire à l'émetteur et au récepteur. Il faut considérer si les gains en robustesse justifient cette complexité.

En résumé, l'implémentation suit le flux de données à travers les différents blocs. La partie cruciale est la conception et l'entraînement de l'auto-encodeur pour qu'il apprenne une représentation robuste aux types de bruit que vous rencontrez dans votre canal OFDM.

IV.9 Conclusion

Dans cette partie nous avons démontré l'efficacité de recours aux techniques du deep learning , notamment les réseaux de CNN à travers un module intelligent appelé auto- encodeur qui permet de modéliser le signal de façon complètement automatique ce qui facilite la compréhension du comportement du signal original et par suite de détecter toute anomalie entachée au signal à travers un bruit impulsif. Cette maîtrise du comportement du signal permet au DAE (Denoising Auto Encoder) de détecter ce bruit et de l'éliminer facilement.

Conclusion générale

Le domaine des télécommunications avance à une vitesse toujours croissante. Aujourd'hui, un utilisateur peut non seulement se connecter à tout moment et en tout lieu, mais il a aussi la possibilité d'accéder à une variété de services via le même réseau, tels que la visioconférence, les jeux vidéo, etc. Les réseaux de nouvelle génération sont conçus en conformité avec le principe de convergence des réseaux. Cette progression incite les chercheurs à perfectionner sans cesse les systèmes de communication en place et à dénicher des solutions plus performantes pour les normes à venir.

Dans ce contexte, notre étude nous a donné l'opportunité de mieux comprendre la récente technologie 5G, notamment la méthode OFDM (Orthogonal Frequency Division Multiplexing). Notre travail se divise en deux sections principales : une section théorique et une seconde centrée sur la simulation dans le but d'obtenir des résultats à discuter.

Le premier volet de notre recherche a englobé une revue littéraire sur le progrès des diverses générations de réseaux de communication, allant de la 1G à la 5G. Par la suite, le second chapitre s'est focalisé sur la technique OFDM. Nous avons minutieusement exposé ses performances, son mécanisme de fonctionnement et clarifié la chaîne de transmission, qui inclut l'émetteur (codage, modulation, modulation OFDM), le canal de transmission (répartition du canal de Rayleigh, bruit AWGN et bruit impulsif) ainsi que le récepteur (démodulation, décodage, etc.). Le troisième chapitre traite des bases théoriques concernant les réseaux de neurones et les auto-encodeurs. La quatrième partie traite des résultats obtenus grâce à la simulation de la chaîne de transmission OFDM effectuée avec le langage de programmation Python dans un environnement de développement Google Colab . Les résultats démontrent l'efficacité de l'auto-encodeur (efficacité) pour éliminer le bruit impulsif.

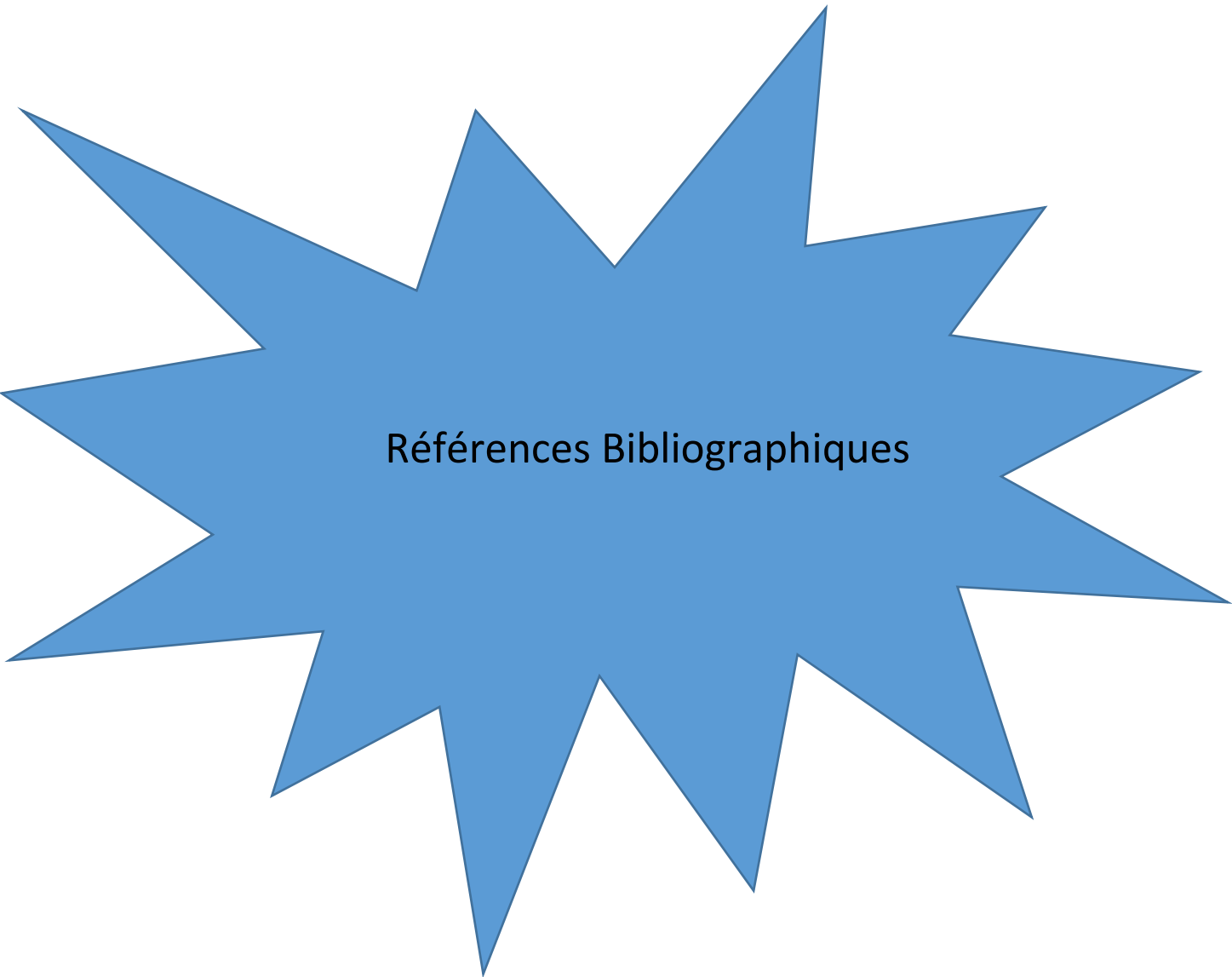
Cette recherche nous a donné une vision du domaine scientifique en relation avec le débruitage des systèmes OFDM. Dans cette modeste démarche, nous avons aussi examiné l'application et la réalisation de l'apprentissage profond (Deep Learning) via des réseaux neuronaux CNN pour réduire les bruits impulsifs, obtenant ainsi des résultats prometteurs.

Néanmoins, Nous avons fait face à des défis lors de notre travail, notamment dans la partie matérielle, où nous étions limités en termes de performances de la machine pour exécuter efficacement les algorithmes. Pour remédier à ce genre de problème nous avons opté pour l'utilisation de l'environnement de développement Google Colab qui nous offrent la possibilité d'utiliser à distance et à titre limité mais efficace ses processeurs ultra puissants, Ainsi, nous avons pu atteindre des résultats de débruitage très positifs, comme le prouve l'évaluation de critères tels que le taux d'erreur binaire (BER) exprimé et expliqué dans le chapitre 04 .

Dans une optique d'avenir, il serait pertinent d'étendre notre recherche à d'autres formes de données telles que les images, les vidéos, etc. Il pourrait aussi être judicieux de réaliser une analyse comparative entre l'autoencodeur et d'autres codes traditionnels.

Les systèmes MIMO (Multiple-Input Multiple-Output) sont des technologies utilisées dans les communications sans fil (comme le Wi-Fi ou la 5G) qui emploient plusieurs antennes à l'émission et à la réception pour améliorer la fiabilité et le débit des transmissions de données.

Dans les publications scientifiques, l'analyse des systèmes MIMO se fait généralement en considérant un bruit de nature gaussienne. Il serait donc judicieux d'examiner ces systèmes sous l'angle de scénarios impliquant du bruit impulsif, voire mixte.



Références Bibliographiques

References bibliographiques

- [1] ECC Report 225, "Establishing criteria for the accuracy and reliability of the caller location information in support of emergency services," tech. rep., CEPT Electronic Communications Committee, Oct. 2014.
- [2] del Peral-Rosado, JA, Raulefs, R., López-Salcedo, JA et Seco-Granados, G. (2017). Etude des méthodes de localisation radio mobile cellulaire : De la 1G à la 5G. *IEEE Communications Surveys & Tutorials* , 20 (2), 1124-1148.
- [3] Scourias, J. (1995) Overview of the global system for mobile communications. University of Waterloo, 4, 7.
- [4] Cristian, T. O. M. A. (2007) Secure Architectures for Mobile Applications. *Informatica Economica*, 11 (4), 89-92.
- [5] Sanchez, J., & Thioune, M. (2013). UMTS. John Wiley & Sons.
- [6] Spiceworks "what is UMTS Universal Mobile Telecommunication System? Meaning, working, Importance, and Application ".2011
- [7] Kuuboore, M., Odai, D. A., & Kotey, A. N. (2023). Telecommunications Wireless Generations: Overview, Technological Differences, Evolutional Triggers, and the Future. *International Journal of Electronics and Telecommunications*, 105-114.
- [8] Cox, C. (2014). An introduction to LTE: LTE, LTE-advanced, SAE, VoLTE and 4G mobile communications. John Wiley & Sons.
- [9] Launay, F., & Pérez, A. (2019). LTE-Advanced Pro: une étape vers le réseau de mobiles 5G. ISTE Group.
- [10] Singh, R., Thompson, M., Mathews, SA, Agbogidi, O., Bhadane, K. et Namuduri, K. (2017, septembre). Stations de base aériennes pour permettre les communications cellulaires en cas d'urgence. En 2017 Conférence internationale sur la vision, l'image et le traitement du signal (ICVISIP) (pp. 103-108) IEEE.
- [11] Ng voice/whaat is an IP Multimedia Subsystem (IMS).12/02/2016
- [12] Agarwal, A. (2017) The 5th Generation Mobile Wireless Networks- Key Concepts, Network Architecture and Challenges (Pubs.sciepub.com) Retrieved from .
- [13] Hong, W., Ko, S., Lee, Y., & Baek, K. (2015). Multi-polarized antenna array configuration for mmWave 5G mobile terminals. 2015 International Workshop on Antenna Technology (iWAT).
- [14] 3GPP TS 38.401: "NG-RAN; Architecture description",version 15.3.0 Release 15 (p10).2014
- [15] Brown, G. (2017). Service-based architecture for 5g core networks. Huawei White Paper, 1.
- [16] 5G Core and EPC interworking variant, Heavy Reading.01/03/2015
- [17] D. L. Goeckel et G. Ananthaswamy, "On the Design of Multidimensionnel Signal Sets for OFDM Systems", *IEEE Transactions on Communications*, Vol. 50, No. 3, pp. 442-452, mars 2002.
- [18] Hanzo, L., & Keller, T. (2007). OFDM et MC-CDMA : une introduction . John Wiley et fils.
- [19] Traverso, S. (2007). Transposition de fréquence et compensation du déséquilibre IQ pour des systèmes multiporteuses sur canal sélectif en fréquence (Doctoral dissertation, Université de Cergy Pontoise).
- [20] Francis Cottet, TRAITEMENT DU SIGNAL SCIENCES SUP Série Aide-mémoire chapitre 4 : Modulation des signaux, pp .57-90.11/09/2003
- [21] Francis Cottet, Traitement des signaux et acquisition de données, chapitre 5 : La modulation , pp 85-117.09/02/2018
- [22] Tesserault, G. (2008) Modélisation multi-fréquences du canal de propagation (Doctoral dissertation, Poitiers).08/05/2006
- [23] Fading channel Models – shodhganga. 2002
- [24] J. D. Parsons. The Mobile Radio Propagation Channel. John Wiley, 2000.
- [25] Bhat, AA et Ahmad, SP (2020). UNE NOUVELLE GÉNÉRALISATION DE LA DISTRIBUTION DE RAYLEIGH : PROPRIÉTÉS ET APPLICATIONS. *Revue pakistanaise des statistiques* , 36 (3).
- [26] Mirza, A., Kabir, S. M., & Sheikh, S. A. (2015). Reduction of impulsive noise in OFDM system using adaptive algorithm. *International Journal of Electronics and Communication Engineering*, 9(6), 1427-1431.

- [27] Chitre, MA, Potter, JR, & Ong, SH (2006). Détection de signal optimale et quasi optimale dans le bruit ambiant dominé par les crevettes claquantes. *IEEE Journal of oceanic engineering* , 31 (2), 497-503.
- [28] Saunders, S. R., & Aragón-Zavala, A. (2007). *Antennas and propagation for wireless communication systems*. John Wiley & Sons.
- [29] Berry, LA (1981). Comprendre la formule canonique de Middleton pour le bruit de classe A. *Transactions IEEE sur la compatibilité électromagnétique* , (4), 337-344.
- [30] Ndo, G., Labeau, F., & Kassouf, M. (2013). Un modèle de Markov-Middleton pour le bruit impulsif en rafale : modélisation et conception du récepteur. *IEEE Transactions on Power Delivery* , 28 (4), 2317-2325.
- [31] Andrei, M., Trifina, L. et Tarniceriu, D. (2015). Capacité du canal de bruit impulsif Middleton de classe A avec entrée binaire. *Mathématiques appliquées et sciences de l'information* , 9 (3), 1291.
- [32] Andrei, M., Trifina, L. et Tarniceriu, D. (2014). Analyse des performances des canaux de relais de décodage et de transfert turbo-codés avec le bruit impulsif Middleton de classe A. *Avancées en génie électrique et informatique* , 14 (4), 35-43.
- [33] En ligne Ghosh, M. (1996). Analyse de l'effet du bruit impulsionnel sur les systèmes QAM multiporteuse et monoporteuse. *Transactions IEEE sur les communications* , 44 (2), 145-147
- [34] Shao, M., & Nikias, CL (1993). Traitement du signal avec moments fractionnaires d'ordre inférieur : processus stables et leurs applications. *Actes de l'IEEE* , 81 (7), 986-1010.
- [35] Farhang-Boroujeny, B. (2013). *Adaptive filters: theory and applications*. John Wiley & Sons.
- [36] Brighta, JM et Killingera, S. (2019). Rectificatif à "Sur la recherche de caractéristiques représentatives des systèmes PV : collecte de données et analyse de l'azimut, de l'inclinaison, de la capacité, du rendement et de l'ombrage du système PV"[Sol. Énergie 173 (2018) 1087-1106]. *Énergie solaire* , 187 , 290-292.
- [37] : Shinde, P. P., & Shah, S. (2018, August). A review of machine learning and deep learning applications. In 2018 Fourth international conference on computing communication control and automation (ICCUBEA) (pp. 1-6). IEEE.
- [38] : https://www.researchgate.net/figure/Schema-dun-neurone_fig1_280792237 , 01.04.2025.
- [39] : Balázs HIDÁSI et al. «Session-based recommendations with recurrent neural networks». In ©2015). eprint: 1511.06939.
- [40] : LONTCHI, C. M. CONTRIBUTION À LA PRÉDICTION DES PERTES DE PUISSANCE SUR UN RÉSEAU ÉLECTRIQUE PAR UN MODÈLE À BASE DE RÉSEAU DE NEURONES. 03/02/2003
- [41] : Mueller, J. P., & Massaron, L. (2019) *Deep learning for dummies*. John Wiley & Sons.
- [42] : Auto-encodeurs en Deep Learning : tout savoir | Blent.ai 11/09/2020
- [43]: A. Felix, S. Cammerer, S. Dörner, J. Hoydis and S. Ten Brink, "OFDM-Autoencoder for End-to End Learning of Communications Systems," 2018 IEEE 19th International Workshop on Signal Processing Advances in Wireless Communications (SPAWC), 2018, pp. 1-5, doi: 10.1109/SPAWC.2018.8445920.